

**ESTABLISHING NORMATIVE REFERENCE DATA FOR ACOUSTIC
VOICE QUALITY INDEX (AVQI) IN TELUGU SPEAKING ADULTS**

KEERTHANA KATTI

Register Number: P01II24S123022

A Dissertation Submitted in Part Fulfillment for the
Degree of Master of Science (Speech -Language Pathology)
University of Mysore



**ALL INDIA INSTITUTE OF SPEECH AND HEARING,
MANASAGANGOTHRI,
MYSORE – 570006
JUNE - 2026**

CERTIFICATE

This is to certify that this dissertation entitled "**Establishing Normative Reference Data for Acoustic Voice Quality Index (AVQI) in Telugu Speaking Adults**" is a bonafide work submitted as part of fulfilment of the degree of Master of Science (Speech Language Pathology) of the student with Registration Number: P01II24S123022. This has been carried out under the guidance of a faculty of this institute and has not been submitted earlier to any other university for the award of any other Diploma or Degree.

Mysuru

8th June, 2026


Dr. M. Pushpavathi

Director

All India Institute of Speech and Hearing

Manasagangothri, Mysuru-570006

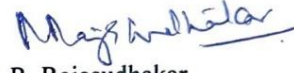
CERTIFICATE

This is to certify that this dissertation entitled "Establishing Normative Reference Data for Acoustic Voice Quality Index (AVQI) in Telugu Speaking Adults" has been prepared under my supervision and guidance. It is also being certified that this dissertation has not been submitted earlier to any other university for the award of any other Diploma or Degree.

Mysuru

June, 2026

Guide



Dr. R. Rajasudhakar

Associate Professor in Speech Sciences

Department of Speech-Language Sciences

All India Institute of Speech and Hearing

Manasagangothri, Mysuru-570006

DECLARATION

This is to certify that this dissertation entitled "**Establishing Normative Reference Data for Acoustic Voice Quality Index (AVQI) in Telugu Speaking Adults**" is the result of my own study under the guidance of Dr. R. Rajasudhakar, Associate Professor in Speech Sciences, All India Institute of Speech and Hearing, Mysuru and has not been submitted earlier to any other university for the award of any other Diploma or Degree.

Mysuru

June, 2026

Registration no. P01II24S123022

ACKNOWLEDGEMENT

**As the heavens are higher than the earth, so are my ways higher than your ways
and my thoughts than your thoughts - Isaiah 55:9**

*I give all glory and thanks to **Almighty God**. What I am today and where I stand today is not by my own wisdom or strength, but solely by his eternal grace, guidance, and endless blessings.*

*I would like to sincerely thank my Guide, **Dr R. Rajsudhakar** (Associate Professor in Speech Sciences) for his calm guidance and encouragement. Having been your student since my bachelor's degree, starting with my Clinical Conference in the third year and continuing through my master's studies, sir, I feel truly fortunate to have learned from your knowledge, wisdom, and mentorship over the years. One of the lessons from you that I will always carry with me is your belief that there can be challenges in teaching, but there is never a bad student who is willing to learn, and I will always be grateful for the opportunity to have been your student.*

*Sincere thanks to **Dr. M. Pushpavathi** for providing me with the opportunity to undertake this research*

*I also extend my heartfelt gratitude to **Dr. Vasanthlakshmi ma'am and Ms sahana mam** , for helping me with the results of my study.*

***Dr. Jesnu Jose Benoy** , Thank you sir, for generously sharing your expertise and helping me in understanding concepts of AVQI. Your timely suggestions and observations helped me in the successful completion of this research. I'm truly grateful for your time and support.*

*Special thanks to **Ms. Susan G. Oommen, Mr. Manas Siddanth, Mr. Santosh Dharamkar,** and **Dr. Rajendra Kumar Porika** for their assistance in recruiting participants and supporting the data collection process.*

*To the two people I owe everything to, **Amma and Nanna**. Thank you for being my constant source of love, strength, and encouragement. Everything I am today is because of your sacrifices, prayers, and faith in me. You always placed my dreams*

before your own needs and never stopped believing in me, even when I doubted myself. This achievement is as much as yours as it is mine. No acknowledgement, however long, can fully capture my gratitude for everything you have done for me.

*To my dearest **Anna**, even though we are miles apart, you have never been distant whenever I needed you. You are always just one call away. Thanks for being my personal ATM and unpaid therapist. You stood by me, when I faced challenges, never judged me and always helped me find my way forward. I know there were many sacrifices and silent struggles behind your support, and I will always be grateful for them. As much as I am thankful to our parents, I am equally thankful to have a brother like you. You will always hold a special place in my heart.*

*Thanks to **Tristan and Namratha** for their smooth coordination in sharing the recorder whenever needed.*

*Thank you, **Subhash**, for your time and effort, and for always thinking of ways to help me with this study. Your help, patience, and belief in me meant more than you think. I am truly grateful for the role you played in this journey*

***Kowshik sir** , Thankyou , for your help in finding participants for this study and for always being just a message away whenever I needed advice regarding administrative matters, grateful for your readiness to help.*

***Siva Sir**, thank you for the conversations, advices, and all the plans you patiently made for me, despite knowing that I might not always follow them and I am very glad to have crossed paths with you.*

*To my darling baddies, **Prasanna, Vidya, Shraddha, Hithyashree, Karunya, Lamu, Deepthi, Vidyashree**, thank you for making this master's journey memorable with your craziness, terrible jokes, unnecessary drama, random gossip sessions, last-minute panic attacks, endless photos and laughter that could probably wake up half the hostel.*

***Thanks to Prasanna and Vidhya**, my technical problem solvers, closest buddies*

***Shraddha**, thank you for the coffees, conversations, and constant companionship*

throughout the Post graduation journey. In a very short time, you came to understand me so well, more than anyone. I'm blessed for your friendship and for always being there for me.

*Thanks to **Hiba K.M., Swathi, and Akshaya**, my wonderful roommates and buddies. Although we did not get much time together to spend as we did in our bachelors , I truly value the friendship, support, and memories we shared along the way.*

*Thanks to my juniors **Haveela, Srinidhi, and Lohith Raghava** the only Telugu people I had around and for contributing samples for this study.*

*My sincere thanks to my church family at **Mysore Hope Centre** for their prayers and support . I am grateful to **Pastor Moses** for his prayers, and to **Brother Shelumiel** and **Sister Udaya** for helping me connect with participants for this study. Heartfelt Thanks to **Sai, Sameena, Kishore, Glory, Dinesh, Mahidhar, Swapna**, and my church family for their participation, friendship, and support throughout this Dissertation Journey. What began as a research project also became an opportunity to build stronger connections and create cherished memories.*

*A note of thanks to a special friend and Master Chef, **Lokesh a.k.a Chinnu**, someone I never expected to meet, yet whose presence became an unexpected blessing towards the end of my master's journey. From helping me find participants to making sure I never stayed hungry. The best thing about you is that you have exactly one solution for every problem "Go eat something" and if that doesn't work, "Sleep and then eat something. Somewhere between all those meals, conversations, and random outings, you became one of the best parts in this Master's journey.*

Lastly, I acknowledge myself for not dropping out every time I said "I can't do this anymore." Against all odds, several breakdowns, and procrastination, I somehow made it to the finish line and while I leave with a master's degree looking back, six years of AIISH gave me more than a degree; it gave me lessons, memories, friendships, and experiences that will stay with me for life.

TABLE OF CONTENTS

Chapter	Title	Page number
	List of Tables	i.
	List of Figures	ii.
I	Introduction	1-7
II	Review of Literature	8-28
III	Method	29-35
IV	Results	36-52
V	Discussion	53-61
VI	Summary and Conclusion	62-69
	References	71-85
	Appendices	86-88

LIST OF TABLES

No.	Title	Page no.
2.1	Summary of AVQI Validation and Normative Studies Across Different Languages and Populations	25
4.1	Mean (M) and standard deviation (SD) values of AVQI in normophonic individuals	37
4.2	Mean (M) and standard deviation (SD) values of the constituent parameters of AVQI among normophonic individuals across Group I, Group II, and Group III	39
4.3	Results of Two-way MANOVA for age comparison on AVQI and its constituent parameters	41
4.4	Results of Two-way MANOVA for gender comparison on AVQI and its constituent parameters	41- 42
4.5	Results of Tukey's HSD Post hoc test for between group comparison of AVQI parameters	45
4.6	Mean (M) and Standard Deviation (SD) values of AVQI and its constituent parameters among dysphonic individuals (group IV)	46
4.7	Results of Mann–Whitney U test comparing normophonic and dysphonic groups	49
4.8	Mean and Standard deviation scores of AVQI and its constituent parameters in native Telugu speaking normophonic adults across three age groups and between males and females	51
4.9	Median and Interquartile range scores of AVQI and its constituent parameters in native Telugu speaking normophonic adults across three age groups and between males and females	52

LIST OF FIGURES

No.	Title	Page No.
3.1	Graphical representation of AVQI results for a normophonic adult	33
3.2	Graphical representation of AVQI results for a dysphonic individual.	34
4.1	Mean AVQI, CPPS, HNR and shimmer Local (%) values between Males and Females across three groups	43
4.2	Mean shimmer Local dB, slope of LTAS, and Tilt through LTAS values between males and females across three groups.	44
4.3 A	Comparison of Mean values of AVQI, CPPS, HNR, Shimmer(%) and Shimmer (dB) among normophonic males, normophonic females and dysphonic individuals	47
4.3 B	Comparison of Mean values of LTAS and Tilt through LTAS among normophonic males, normophonic females and dysphonic individuals	48

CHAPTER I

INTRODUCTION

Voice is crucial for daily communication. The speaking voice inherently communicates information about the individual, with voice quality serving as a key means by which speakers convey their psychological, physical and social traits to others (Laver, 1980). The human voice is produced by vibration caused by the flow of the vocal folds and requires a complex interaction between acoustics, tissue motion, and flow dynamics (Thomson, 2024). Voice is determined as normal when the voice quality is clear; the pitch and the loudness of the voice are appropriate for age, gender and situation; the voice is produced without excess effort, pain, strain, or fatigue; and the voice is satisfactory to the speaker in meeting his or her occupational, social and emotional vocal needs. When an individual's voice differs significantly from normative data with respect to quality, pitch or loudness this suggests that he/she may have dysphonia (Coyle, Weinrich & Stemple, 2001). The term dysphonia thus encompasses a wide range of voice changes that may arise from structural, neurological, functional, or psychogenic causes, all of which can disrupt the typical vibratory pattern of the vocal folds and the resulting acoustic output.

The study of voice is vital not solely because of its biological and linguistic functions, but also because voice abnormalities are prevalent in people around the world. A landmark population study discovered that the prevalence across the lifespan of voice abnormalities was around 29.9%, with 6.6% of studied individuals reporting a current voice disorder (Roy et al., 2005). This highlights the need for reliable tools to evaluate and monitor voice quality. A Speech Language Pathologist (SLP) who deals with dysphonic patients must do a comprehensive evaluation in order to determine the accurate diagnosis, which later serve as the basis for management.

Across the studies, voice quality has often been evaluated using various perceptual and objective measures. Perceptual measurement of the voice is a method that employs auditory-perceptual processes, together with incorporation of inputs from a rating system to make effective judgments on the typical nature and appropriateness of an individual's voice (Aronson & Bless, 2009). However, this auditory perceptual voice assessment is not considered as an objective measurement because it utilizes the listener's ears to interpret what they are listening. Additionally, listener's experience and training influence auditory perceptual evaluation of voice. Due to its subjective nature, it is further influenced by the listener's bias. Despite all the disadvantages, it is always considered the gold standard because it provides information readily to the examiner and is found to be economical in its usage (Oates, 2009). The advantage of perceptual analysis is that it is easily accessible for SLPs and laryngologists for daily use in their clinical settings. The most widely employed perceptual evaluation assessment tools are GRBAS scale (Hirano, 1981), Consensus Auditory Perceptual Evaluation of Voice (CAPE-V) (Kempster et al., 2009), Buffalo Voice III Profile (Wilson, 1987), Vocal Profile Analysis Scheme (Laver et al., 1981), and so on. Furthermore, self-rating measures such as the voice handicap index are used to quantify the impact of vocal problems on individual's daily activities and quality of life. Voice is also evaluated objectively through acoustic, aerodynamic and imaging techniques.

Regardless of the significance of perceptual measures, objective evaluation is crucial in terms of its reliability, describing the voice parameters more quantitatively and serves as a diagnostic indicator. Acoustic measures, electroglottographic techniques, aerodynamic measures, and so on are some of the objective measurements used for evaluating voice parameters. However, there is an increase in evidence of using the acoustic parameters of voice assessment to evaluate the vocal function due to its non-invasive method of obtaining the data and as it

focuses directly on the output characteristics. Acoustic analysis serves two purposes. During evaluation, it provides indirect evidence of the severity of voice problem and at follow-up, it helps to evaluate effects of rehabilitation plan (Scherer, Vail, & Guo, 1995).

Acoustic analysis of voice includes various spectral and cepstral parameters which can be used for both diagnoses as well as tracking intervention efficacy in voice disorders. These are frequently used by the Speech-language pathologist as it is non-invasive, time-saving, and easy to interpret. It has been used for diagnostic investigations as well as to track the treatment efficacy (Carding et al., 2009). It involves procedures such as inverse filtering, auto correlation, spectrum, cepstrum to extract the frequency related measures, amplitude related measures, perturbation related measures, noise/harmonic related measures, and measure of voice continuity (Hirano et al., 1988; Wolfe et al., 1995; Dejonckere & Lebacqz, 1996 and Picirillo et al., 1998). However, Awan and Roy (2006) stated that single parametric measurement such as cepstral and spectral based measures have limited validity for diagnostic outcomes. Although a majority of the acoustical parameters were easy to analyze and interpret, they were shown to have limited test-retest reliability and poor correlation with perceptual analysis (Bauser & Drinnan, 2011; Karnell, Hall, & Landanl, 1995; Bough et al., 1996). This has led to a surge in the trend of using a multivariate objective analysis of voice rather than unidimensional analysis.

Some of the recently introduced multi-parametric acoustic measurement to identify the severity of dysphonia in voice are Dysphonia Severity Index (DSI) and Acoustic Voice Quality Index (AVQI). These parameters were reported to be effective when they used in weighted combinations such as Dysphonia Severity Index (Wuyts et al., 2000), Acoustic Voice Quality Index (Maryn, Bodt & Roy, 2010), Cepstral spectral index of dysphonia (Peterson et al., 2013).

The Dysphonia severity index (DSI) (Wuyts et al., 2000) is a multiparametric measurement reported to be a robust measure in different studies (Timmermans et al., 2004). DSI considers Maximum Phonation Time (MPT), highest frequency, lowest intensity and jitter to reach a numerical value that reflects the voice quality of a given individual. In a study by Hakkesteegt et al. (2008), it was revealed that the mean DSI scores could be highly stable across subject groups were well correlated with the GRBAS perceptual scale. In the Indian context, Jayakumar and Savithri (2012) developed the normative for DSI and compared it with the European population in terms of Highest F0, MPT as well as the DSI values. One of the drawbacks of DSI is that connected speech is not taken into consideration. To improve ecological validity, it is essential to analyze both sustained phonation as well as connected speech, as this allows for a more realistic representation of everyday voice use.

Maryn et al. (2010) investigated the viability and diagnostic accuracy of combining continuous phonation with speech in the assessment of the quality of voice using -Acoustic Voice Quality Index (AVQI). The six parameters assessed in AVQI are Smoothed Cepstral Peak Prominence (CPPS), Harmonics-to-Noise Ratio (HNR), Shimmer Local (SL), Shimmer local dB (ShdB), slope of long-term average spectrum (slope) and tilt of the trendline through the long-term average spectrum (tilt). AVQI possesses a script designed for the automated assessment of dysphonia severity, utilizing an analysis of a combined sample involving sustained vowel and connected speech. Hence, AVQI is constructed with the following formula: “AVQI = $2.571 \cdot (3.295 - 0.111 \cdot \text{CPPS} - 0.073 \cdot \text{HNR} - 0.213 \cdot \text{SL} + 2.789 \cdot \text{ShdB} - 0.032 \cdot \text{Slope} + 0.077 \cdot \text{Tilt})$ ” (Maryn et al., 2010). This script is exclusively compatible with Praat software. A score of 2.95 or below obtained on AVQI identifies the sample to be normophonic adults (Maryn et al., 2010). The AVQI score, ranges from 0 to 10, serves as an indicator of voice severity, with 0 representing normal voice and 10

signifying profoundly abnormal voice on the severity continuum (Maryn et al., 2010).

A high positive correlation of approximately 0.78 was identified between AVQI and the G parameter of GRBAS. This indicates that as the AVQI score increases, reflecting a higher degree of hoarseness, there is a corresponding deterioration in overall voice quality, and conversely, a lower AVQI score is associated with better voice quality (Maryn et al., 2010).

Version 03.01 of AVQI has undergone significant modifications, particularly in the adjustment of parameter weights, leading to its current form. AVQI holds good diagnostic accuracy and cross-linguistic validity (Maryn et al., 2014). The script for obtaining AVQI (v03.01) reported by Barsties and Maryn (2015b), contains the formula, “ $AVQI = (4.152 - (0.1777 * CPPs) - (0.006 * HNR) - (0.037 * Shim) + (0.941 * ShdB) + (0.01 * Slope) + (0.093 * Tilt) * 2.8902$ ”.

Jayakumar and Benoy (2022) reviewed eight studies on AVQI (v03.01) and seven studies on AVQI (v02.02) and reported an overall sensitivity and specificity of 0.85 and 0.92 for AVQI (v02.02) and 0.82 and 0.92 for AVQI (v03.01), respectively. Additionally, it has been found that the area under the curve (AUC) was found slightly better for AVQI (v03.01)(0.94) than AVQI (v02.02) (0.92). The diagnostic threshold of AVQI (v02.02) is within the range of 2.72 to 3.33 while, in case of AVQI (v03.01), it is found to be in the range of 1.33 to 3.15 (Jayakumar & Benoy, 2022).

There is no effect of gender on AVQI (v02.02) as reported by Latoszek et al. (2019). Shabnam and Pushpavathi (2021) in AVQI and DSI in the Indian context even though there was significant effect seen in their individual parameters. Jayakumar et al. (2022) also reported that there is no effect of gender on AVQI (v02.02), however, their study across the lifespan revealed that age was a factor on AVQI (v02.02).

There are numerous studies that reports AVQI values in adult population. In

Indian context, Benoy (2017) has developed and validated an AVQI (v02.02) reference data with perceptual measurements between healthy normal individuals and individuals with dysphonia. Vishali (2019) has established standard reference data for AVQI (v02.02) in native Tamil speakers between the age range of 20 to 50 years. Seshasri (2018) has established normative data for AVQI (v02.02) for Kannada speaking typically developing children in the age range of 10 -12 years. Varna (2024) has established the normative data for AVQI v03.01 in Malayalam and Kannada speaking typically developing children in the age range of 5- 8 years. To date, no study has established normative reference data for AVQI v2.02 in Telugu-speaking adults. Therefore, the present study focuses on establishing normative reference data for AVQI v2.02 in healthy normal Telugu speaking adults.

1.1 Need for the study

The Acoustic Voice Quality Index (AVQI) is an objective multiparametric measure widely recognized for evaluating voice quality. Since its development, it has been validated in several languages across the world like Dutch, German, English, French, and Finnish (Maryn et al., 2014), the Altaic language Korean (Maryn et al., 2016; Kim et al., 2018), and the Indo-European language Lithuanian (Uloza et al., 2016) and Persian (Zainae et al., 2022), where AVQI has consistently demonstrated its clinical utility in distinguishing normal from dysphonic voices and correlate with perceptual judgments.

Telugu is one of the Dravidian languages and the fourth most spoken native language in India, with approximately 81 million speakers, accounting for 6.7% of the country's population (Census of India, 2011). It is the native language spoken by the residents of the South Indian states of Andhra Pradesh and Telangana. The phonetic system of Telugu is extensive, featuring 12 vowels and 36 consonants, with clear vowel length distinctions and a range of phonological transformations largely influenced by sandhi rules that affect sound changes at word boundaries during connected speech (Bhaskararao & Ray, 2017). Despite these rich and complex

phonetic characteristics of Telugu, the AVQI has not been investigated so far. At present, there is no normative reference AVQI data for Telugu-speaking adults, and no validation studies with Telugu speakers who suffer with dysphonia. In addition, although research in other Indian languages has examined the role of age and gender, no such studies have been conducted in the Telugu population. Therefore, there is a need to develop clinical reference data on AVQI in the normophonic (healthy) Telugu-speaking population, examine the influence of age and gender on AVQI and its constituent parameters, and explore the applicability of AVQI in Telugu-speaking individuals.

1.2 Aim of the study

The present study is aimed to establish the normative reference data for Acoustic Voice Quality Index (AVQI) in native Telugu-speaking normo-phonic adults within the age range of 20 to 65 years.

1.3 Objectives of the study

- i. To determine normative reference data for the Acoustic Voice Quality Index (AVQI) in native Telugu-speaking healthy young adults aged 20 to 35 years; middle aged adults aged 36 to 50 years and older adults aged 51 to 65 years.
- ii. To determine the effect of age on AVQI scores and its constituent parameters in the study population.
- iii. To examine the influence of gender on AVQI scores and its constituent parameters in Telugu speakers.
- iv. To validate the Acoustic Voice Quality Index (AVQI) by testing it in individuals with dysphonia who are native Telugu speakers.

CHAPTER II

REVIEW OF LITERATURE

The human voice is the result of the coordinated activity of several physiological systems, involving the interaction of respiratory, phonatory, and resonatory subsystems rather than a single isolated process. The coordinated interaction of airflow through the glottis, vocal fold vibration, and acoustic wave transmission within the vocal tract mechanically produces voice, which reveals the close relationship between physiological activity and physical processes (Calvache et al., 2023). Physiologically voice production can be understood at three levels: the subglottal system, which is responsible for controlling airflow and pressure; the glottal level, where vocal fold vibration occurs; and the supraglottal system, where sound is modified through vocal tract resonance. Voice is usually defined in clinical and perceptual terms through attributes like pitch, loudness, and voice quality. These attributes are perceived collectively rather than as separate elements, contributing to the listener's overall impression of a speaker's voice. The variations in these characteristics arise from the interaction of aerodynamic factors, vocal fold behaviour, and acoustic resonance, contributing to the uniqueness of each human voice. Based on the results of recent clinical studies, it is essential to identify any variations in the voice quality in order to distinguish between pathological and normal variation of voice and for guiding evidence-based voice assessment and intervention (Roy et al., 2010; Lechien et al., 2023).

The term dysphonia is used for any kind of change in voice quality in

terms of pitch, loudness, or any loss of efficiency of vocal production (American Speech-Language-Hearing Association [ASHA], 2015. Epidemiological data from recent times suggest that voice issues are much more common in all populations. As per more recent statistics from studies conducted in 2024, roughly 12.6% of the general population reports current vocal issues, while 20.6% report having voice disorders at some point in their lives (Roy et al., 2020). Research by Fujiki and Thibeault (2024) indicates that voice disorders affect a considerable number of adolescents, with 24.3% reporting past experiences and 7.4% reporting current voice problems. This suggests that voice disorders affect people early in life and may last a lifetime. Due to this widespread prevalence and variability along with the multidimensional nature of voice production, it becomes extremely essential for accurate identification and systematic evaluation of voice disorders. Therefore, assessment and diagnosis of voice disorders require comprehensive, multidimensional evaluation approaches. A Speech-Language Pathologist (SLP) dealing with dysphonic patients must conduct a comprehensive evaluation in order to determine an accurate diagnosis, which subsequently serves as the basis for evidence-based management.

2.1 Assessment of Voice

Assessment of voice is a critical component in the diagnosis and management of voice disorders. Clinical voice evaluations often involve auditory-perceptual evaluation, instrumental acoustic analysis, aerodynamic evaluation, laryngeal examination, and patient-reported outcome measures. Among these components, auditory-perceptual evaluation is seen to be essential because it reflects how listeners perceive voice abnormalities in

everyday communication contexts (Oates, 2009; Kempster et al., 2009). Instrumental techniques, especially acoustic analysis, are used to supplement perceptual assessments by offering objective, measurable data about the properties of voice signals (Patel et al., 2018).

2.2 Subjective Evaluation of Voice

Subjective evaluation of voice is also known as psychoacoustic evaluation of voice because it relies mainly on listener's hearing. It is regarded as the 'gold standard' method for the clinical assessment of voice and speech. The main advantages of perceptual evaluation are its easy availability and it does not require any complex instrumentation for its utilization in clinical settings (Nerurkar, 2017). Despite its widespread use, it has been strongly criticized for its subjective nature. The main limitation of perceptual evaluation is that it is primarily based on listener's hearing and is subject to variability due to factors such as the listener's hearing threshold and experience; type of samples used such as phonation, reading, spontaneous speech, and conversation; and subject-related factors (Nemr et al., 2005).

A variety of perceptual rating measures have been developed to overcome these constraints and improve listener agreement in judging voice quality. These perceptual scales make use of visual analog scale, equal appearing interval scale, Likert scale etc for making the judgement regarding voice quality. In literature, most commonly used scales for perceptual analysis of voice are Grade, Roughness, Breathiness, Asthenia, and Strain (GRBAS) scales (Hirano, 1981), Vocal Profile Analysis Scheme (Laver et al., 1981), Consensus Auditory Perceptual Evaluation of Voice (CAPE-V) (Kempster, Gerratt, Verdolini, Barkmerer, Kraemer & Hillman, 2009), and Buffalo Voice

III Profile (Wilson, 1987).

2.2.1 GRBAS Scale

The 'GRBAS Scale' was introduced by Hirano, (1981). It uses a four-point rating scale for all five characteristics, ranging from 0 (normal) to 3 (extreme), and is intended as a baseline assessment of voice quality. Every measure of the scale denotes a different aspect of phonation: R stands for roughness, B for breathiness, A for asthenic (weakness), S for strain, and G (Grade) for the overall severity of the voice (Hirano, 1981). Among these parameters, the *Grade* component is most widely used as a global indicator of dysphonia severity due to its ability to summarize overall perceptual abnormality. Owing to its simplicity, minimal time requirement, and feasibility in routine clinical settings, GRBAS rating scale continues to be one of the most commonly adopted perceptual voice assessment tools in both clinical and research contexts (Oates, 2009; Kempster et al., 2009).

Limitations of the GRBAS Scale

Even as it is extensively used, it is important to understand that GRBAS is fundamentally a subjective tool whose reliability depends on the expertise of the listener, leading to variability across clinicians. Literature reports that inter-rater reliability is generally lower than intra-rater reliability, with parameters such as *Strain* and *Asthenia* showing greater inconsistency than the *Grade* parameter (Vaz Freitas et al., 2014). Moreover, since GRBAS uses a four-level ordinal scale, it lacks sensitivity to detect minor changes in voice quality when compared to continuous rating scales (Nemr et al., 2012). Lastly, GRBAS focuses mainly on auditory-perceptual characteristics of voice quality while failing to address the psychological implications of voice

problems, as witnessed by low correlations with patient-reported outcomes (Vallur & Jagadeeswaran, 2025). These limitations support the use of GRBAS as a complementary rather than a standalone measure in voice assessment.

2.2.2 Vocal Profile Analysis (VPA) scheme

The Vocal Profile Analysis (VPA) scheme, first introduced by Laver (1980, 1991), is a scheme of analysis using a descriptive auditory-perceptual method that utilizes trained listeners to analyze the voice quality based on physiological changes to both laryngeal and supralaryngeal systems during speaking or reading activities. This scheme evaluates voice characteristics in comparison to a predetermined neutral voice sample and numerically represents the deviation for phonatory and articulatory features. The VPA scheme evaluates several aspects of voice production, such as prosodic characteristics like pitch and loudness, phonation characteristics, vocal tract characteristics, and overall muscular tension. A six-point ordinal scale is used to rate deviations from the neutral voice setting, with degrees 1-3 denoting moderate deviations and degrees 4-6 denoting extreme deviations (Laver, 1991; Beck, 2018). While VPA allows for precise, multidimensional descriptions of habitual voice quality (Oates, 2009), there are issues with using this analysis as a clinical tool due to its poor reliability. Webb et al. (2004) reported poor to fair inter-rater reliability for VPA ($K = 0.00\text{--}0.29$) even when ratings were performed by experienced clinicians, which they attributed to the complexity of the scheme. The lack of practical application of VPA was due to both the high degree of training needed for the tool and the variability among raters.

2.2.3 Buffalo Voice Profile rating scale

The Buffalo Voice Profile (Wilson, 1987) is frequently used in rating

voice problems and as a guideline for voice therapy. It contains seven equal-appearing intervals with ‘1’ meaning a slight deviation and ‘7’ a severe deviation. The Profile consists of 12 key dimensions including: laryngeal tone, laryngeal tension, vocal abuse, loudness, pitch, vocal inflections, pitch breaks, diplophonia, resonance, nasal emission, rate, and overall voice efficiency. The rater circles the appropriate term listed under each item. The inclusion of parameters which are not related to the voice quality such as (vocal strain, intermittent aphonia) of an individual is the major limitation of using this scale for evaluation (Carding et al., 2000).

2.2.4 Consensus Auditory-Perceptual Evaluation of Voice (CAPE-V)

The Consensus Auditory-Perceptual Evaluation of Voice (CAPE-V) is a standardized protocol that uses visual analogue scales to assess the voice quality perceptually (Kempster et al., 2009). The tool aims at assessing the overall degree of severity, as well as various vocal aspects, including roughness, breathiness, strain, pitch, and loudness, across sustained vowels and connected speech. There have been a number of findings that showed good intra-rater and inter-rater reliability, particularly for overall severity, when trained listeners follow standardized procedures (Zraick et al., 2011; Nemr et al., 2012). High correlation between the overall severity score of CAPE-V and the score of GRBAS grade “G” further confirm the concurrent validity of the two perceptual assesment tools. High correlations between CAPE-V overall severity and GRBAS Grade “G” further support the concurrent validity of the two perceptual assessment tools. Nemr et al. (2012) found high levels of reliability and agreement between GRBAS and CAPE-V ratings, and Liu et al. (2025) found strong correlations between the two scales corresponding

parameters, particularly CAPE-V Overall Severity and GRBAS Grade ($r_s = 0.91$), showing consistent perceptual interpretation across instruments. However, CAPE-V remains a subjective, listener-dependent measure, the variability across raters as well as the need for training and calibration have been identified as key drawbacks, particularly in routine clinical settings. It is therefore advised to use CAPE-V in conjunction with objective acoustic measures as part of a multidimensional assessment.

Overall, perceptual voice assessment tools such as GRBAS, CAPE-V, VPA, and the Buffalo Voice Profile provide valuable clinical insight into voice quality as perceived by listeners. However, their subjective nature and variability across raters limit consistency and quantification of voice severity. These limitations highlight the need for objective voice evaluation methods that offer reliable and reproducible measures to complement perceptual assessment.

2.3 Objective evaluation of voice

Objective evaluation of voice is one of the best assessment methods for clinical voice evaluation. Objective evaluation of voice involves the use of instrumental techniques to quantitatively assess vocal function by measuring specific acoustic, aerodynamic, and physiological parameters. Unlike perceptual evaluation, which relies on listener judgment and is susceptible to inter- and intra-rater variability, objective assessment provides numerical and reproducible data that help substantiate perceptual findings as part of a comprehensive voice evaluation. Objective measures are essential for confirming the presence of voice disorders, documenting baseline voice characteristics, monitoring treatment outcomes, and conducting clinical research due to their consistency and reduced subjectivity (Yu, Ouaknine,

Revis, & Giovanni, 2001; De Bodt, Barsties, & Maryn, 2010).

Objective voice evaluation methods may be broadly categorized into invasive and non-invasive techniques. The invasive methods include laryngoscopy and videostroboscopy, which provide visual information regarding vocal fold structure and vibratory patterns and aid in understanding the physiological basis of voice disorders. The non-invasive methods include acoustic analysis of voice aerodynamic analysis of voice and electroglottography. Among these methods, acoustic voice analysis is commonly preferred because it is easy to administer, non-invasive, and provides objective and quantitative data that support perceptual judgments (Awan & Roy, 2015).

Acoustic analysis of voice offers an objective way to measure vocal parameters related to frequency, amplitude, perturbation, and noise components. Measures related to frequency, such as fundamental frequency and pitch range, help assess vocal pitch characteristics, while amplitude measures like habitual intensity, reflect vocal loudness. Perturbation measures such as jitter and shimmer, evaluate the stability and periodicity of vocal fold vibration, and noise-related measures, such as harmonic-to-noise ratio, indicate the level of turbulence in voice. These measures have been widely used in both clinical and research settings to assist perceptual voice assessment (Yu et al., 2001).

Early investigations aimed to evaluate whether acoustic parameters can accurately measure the way the voice quality is perceived by the listeners. That is, to explore any relations between subjective and objective measurement of voice quality. Indeed, those investigations demonstrated varying correlation

coefficients and the fact that one acoustic measure cannot reliably represent overall perceptual judgments of voice quality.

According to Kreiman et al. (1994), the overall perceptual impression of voice quality cannot be reliably assessed by using an isolated acoustic parameter indicating multidimensional nature of voice. Moreover, significant correlations between the Grade component of the GRBAS scale and some acoustic measures such as noise-to-harmonic ratio and voice turbulence index have been identified by Bhuta, Patrick, and Garnett (2004); nevertheless, these relation cannot be reliably applied for all acoustic measures. These limitations with single-parameter acoustic analysis were confirmed by additional studies. Bauser and Drinnan (2011) found significant variation in the correlations between commonly used acoustic measures and perceptual severity ratings, indicating low validity and reliability when such measures are utilized independently. The above findings indicate that though objective acoustic measures offer significant quantitative information, overall perceptual experience of dysphonia is not accurately represented by individual characteristics.

In response, researchers began to support the use of multiparametric approaches for objective voice evaluation. In a meta-analysis study conducted by Maryn et al. (2009), it was shown that a combination of spectral and cepstral features demonstrates greater consistency with perceptual voice quality assessment than individual parameters, thus proving that perceptual voice quality emerges as a consequence of several acoustic characteristics changing simultaneously. Similarly, supporting this view, Yu, Pang, and Dong (2014) noted that dysphonic voices demonstrate simultaneous changes in frequency,

amplitude, and noise-related parameters, hence highlighting that combining multiple parameters will yield better clinical utility.

Later studies focused on finding more reliable ways to measure the severity of dysphonia by combining multiple acoustic parameters instead of relying on a single measure alone. Compared to conventional single acoustic parameters, studies conducted during this phase showed that cepstral and spectral measures are more closely correlated with perceptual severity ratings and remain more consistent across different voice types and recording conditions. (Uloza et al., 2018; Maryn, De Bodt, & Roy, 2019). This shift in evidence resulted in the development of composite acoustic indices, which combine multiple objective parameters into a single score to overcome the limits of both perceptual evaluation and single-parameter acoustic analysis.

Literature reports few such multiparametric measures of voice quality including Dysphonia Severity Index (DSI) (Wuyts et al., 2000), Cepstral Spectral Index of Dysphonia (CSID) (Peterson et al., 2013) and Acoustic Voice Quality Index (AVQI) (Maryn et al., 2010).

2.3.1 Dysphonia Severity Index (DSI)

Wuyts et al. (2000) developed the dysphonia severity index from multivariate analysis of 387 subjects with the goal to describe voice quality within objective terms after instrumental analysis using four out of 13 parameters from both aerodynamic and acoustic perturbation measures. Among these 13 measures, only four voice measures were statistically selected using a stepwise logistic regression procedure, representing the minimal set of instrumental measures that could best predict perceptual severity. These four

measures were jitter percent, maximum phonation time of sustained vowel /a/, the highest frequency value and the minimum intensity level. From the study, the DSI was constructed as: $DSI = 0.13 \times MPT + 0.0053 \times F0\text{-High} - 0.26 \times I\text{-low} - 1.18 \times \text{Jitter } (\%) + 12.4$. They compared the perceptual means of “G” from the GRBAS scale, in order to reflect the properties of DSI. They found a linear relationship between the subjective Grading and the more objective dysphonia severity index, indicating that a voice is perceptually rates as hoarseness, the more negative its DSI value becomes. Hence, DSI is constructed so that a perceptually normal voice corresponds with a DSI of + 5 (positive sign) and a severely dysphonic voice corresponds with a -5 (negative sign) , but score beyond this range are also possible (higher than +5 or lower than -5). An advantage of DSI is that the parameters can be obtained relatively quick and easy by speech language pathologists in daily practice.

Several studies reported that DSI is a good correlate of the perceptual dysphonia severity (Wuyts et al., 2000; Hakkesteegt et al., 2006; Neelanjana, 2011).

2.3.2 Cepstral Spectral Index of Dysphonia (CSID)

The Cepstral Spectral Index of Dysphonia (CSID) (Peterson, Roy, Awan, Merrill, Banks & Tanner, 2013) is a multivariate assessment of dysphonia severity that considers both sustained vowels and connected speech. CSID includes several components derived independently from sustained vowel and connected speech samples, such as cepstral and spectral parameters including Cepstral Peak Prominence (CPP), the low-to-high spectral energy ratio (L/H spectral ratio), and its standard deviation. It automatically computes the severity of dysphonia for each task. The CSID value varies between 0 and

100, and occasionally it is less than or larger than what is generated, indicating a voice that is profoundly aperiodic or highly periodic, respectively.

Although the CSID is a multiparametric measure, its results may be influenced by recording conditions and it cannot fully replace auditory-perceptual evaluation because it measures acoustic characteristics of the voice signal, including regularity, periodicity, and spectral structure, rather than the listener's perception of voice qualities such as roughness, breathiness, strain, and overall dysphonia severity (Peterson et al., 2013; Maryn et al., 2014).

2.3.4 Acoustic Voice Quality Index (AVQI)

Maryn et al. (2010) introduced a novel method of objective voice assessment that combines sustained vowel phonation and connected speech into a single acoustic analysis. The Acoustic Voice Quality Index (AVQI) was developed as a result of this work. The authors intended to develop a measure that could yield a single numerical score indicating overall dysphonia severity while maintaining a significant correlation with auditory-perceptual judgments.

It is a multivariate model integrated into the free software program, Praat. It constitutes six acoustic parameters such as the Smoothed Cepstral Peak Prominence (CPPS), Harmonics-to-Noise Ratio (HNR), Shimmer Local (SL), Shimmer local dB (ShdB), slope of long-term average spectrum (slope) and tilt of the trendline through the long-term average spectrum (tilt). Hence, AVQI is constructed with the following formula: “AVQI = $2.571 \cdot (3.295 - 0.111 \cdot \text{CPPS} - 0.073 \cdot \text{HNR} - 0.213 \cdot \text{SL} + 2.789 \cdot \text{ShdB} - 0.032 \cdot \text{slope} + 0.077 \cdot \text{tilt})$ ” (Maryn et al., 2010).

Across the timeline, AVQI version was revised to maximize its feasibility and to reduce the processing steps (Maryn & Weenink, 2015).

AVQI version alpha initially was calculated using the computerized program called *Speech Tool* (James Hillenbrand, Western Michigan University, Kalamazoo, MI, USA, 2011) to derive the CPPS parameter and Praat software (Boersma & Weenink, version 4.6.15, <https://www.praat.org/>) is used to derive the acoustic measurements such as HNR, Shimer local, Shimmer dB, slop and tilt (Maryn et al., 2010). Further, the AVQI beta version (second version) was developed in such a way that the complete derivation of values including CPPS can be obtained in Praat software alone. Further, the results of both the original version and second version yielded highly acceptable approximation (Maryn et al., 2015).

AVQI version 3 was developed with a notion to establish an equal proportion of sustained phonation of vowel and continuous speech sample to reach increased ecological validity. As the length of speech sample was much smaller for the analysis after the separation of voice and voiceless segments, the length of continuous speech was increased from 17 to 22 syllables to roughly 34 syllables (Barsties & Maryn, 2015). As AVQI evolved over time, subsequent studies compared the diagnostic performance of various AVQI versions across languages and clinical populations.

2.3.5 AVQI version 2 Vs version 3

Several validation studies across languages have compared the diagnostic performance of AVQI v02.02 and v03.01. A systematic review and meta-analysis by Jayakumar and Benoy (2022) reported that both versions demonstrate high diagnostic accuracy, with pooled sensitivity and specificity of 0.85 and 0.92 for v02.02, and 0.82 and 0.92 for v03.01,

respectively. Although the Area Under the Curve was slightly higher for v03.01 (0.94) compared to v02.02 (0.92), the difference was marginal, indicating that both versions are clinically robust.

Notably, the diagnostic cutoff values for AVQI were found to differ across versions and languages. For AVQI v02.02, the reported threshold values generally ranged from 2.72 to 3.33, while for AVQI v03.01, the cutoff values ranged from 1.33 to 3.15 (Jayakumar & Benoy, 2022; Latoszek et al., 2020). Studies suggest that these variations are related to language-specific phonetic characteristics and the number of syllables used in the connected speech sample, particularly in AVQI v03.01 (Englert et al., 2021).

Despite these methodological improvements in AVQI v03.01, AVQI v02.02 is still widely used in research and clinical practice across different languages. The continued use of both versions highlights the need for separate clinical reference values based on the specific AVQI version and language to ensure accurate clinical interpretation.

2.4 WESTERN STUDIES ON AVQI

Uloza et al. (2018) explored the use of AVQI version 02.02 in a Lithuanian-speaking population. The study included 60 participants, consisting of 30 individuals with voice disorders and 30 individuals with healthy voices. For the analysis, recordings of connected speech and sustained phonation of vowel /a/ were collected. A stable mid-portion of the sustained vowel selected and voiced segments were extracted from the connected speech sample. Using the GRBAS scale's “ Grade ” parameter, experienced clinicians performed perceptual evaluation. The authors reported a strong correlation between AVQI

scores and perceptual severity ratings. Receiver Operating Characteristic (ROC) analysis revealed an AUC of 0.93, indicating high diagnostic accuracy. The optimal AVQI cut-off value was identified as 2.97. The authors concluded that AVQI v02.02 is a reliable and valid objective measure for assessing voice quality in Lithuanian speakers and emphasized the importance of language-specific threshold values.

Kim et al. (2018 a) examined its diagnostic accuracy in a Korean-speaking population. The study included 1,524 participants, comprising 113 normophonic individuals and 1,411 individuals with voice disorders. Sustained vowel /a/ phonation and reading of the Korean “Walk” passage were recorded, from which a 2-second mid-portion of the vowel and 26 syllables of connected speech were extracted for analysis. Auditory-perceptual evaluation was performed by five speech-language pathologists using the Grade parameter of GRBAS rating scale and overall severity from CAPE-V. Receiver operating characteristic analysis identified an AVQI cut-off value of 3.33, with area under the curve values between 0.970 and 0.977, indicating excellent diagnostic accuracy.

The validity of AVQI version 02.06 and version 03.01 among Persian speakers was tested by Zainae et al. (2022). The study involved 180 participants, with 100 being normophonics and 80 subjects with voice disorders. Voice samples were obtained using sustained voicing of vowel /a/ and connected speech based on standardized Persian sentences and concatenated after which they could be analyzed acoustically. Five experienced speech-language pathologists were involved in auditory-perceptual evaluation, in which the ‘ Grade ’ component of the GRBAS scale and overall severity scores were

assessed by using the CAPE-V protocol. Receiver operating characteristic analysis revealed AVQI v02.06 cut-off values ranges between 3.47 and 4.04, yielding sensitivities ranging from 85.5% to 88.3% and specificities ranging between 79.1% and 85.6%, respectively. For AVQI v03.01, cut-off values of 3.03 and 3.07 demonstrated sensitivities ranging from 74.5% to 85.5% and specificities between 91.4% and 97.7%. The authors found that there is a high concurrent validity and a high diagnostic accuracy in both versions of AVQI and hence the authors affirmed that they can serve as objective means of evaluating severe dysphonia across populations who speak Persian language.

Fantini et al. (2023) validated AVQI version 03.01 in Italian-speaking adults. The study included 150 participants, with 100 dysphonic patients and 50 Normophonics. Voice recordings comprised sustained vowel phonation and connected speech, with the speech sample standardized using a language-specific standardized syllable number (SSN=25 syllables). The SSN was determined based on the syllabic characteristics of the Italian language allowing a fixed and comparable amount of voiced connected speech for analysis across speakers. The study reported a cut-off value of 2.35 for AVQI version 03.01, which yielded a sensitivity of 90% and a specificity of 92% in distinguishing dysphonics from normophonics. Subsequent to its validation across various International languages, AVQI was further investigated within the Indian linguistic context.

2.5 INDIAN STUDIES ON AVQI

In the Indian context, Benoy (2017) developed and validated AVQI v02.02 reference data by comparing vocally healthy individuals with

individuals presenting dysphonia. The study involved 120 adults of 20-50 years of age, comprising 100 vocally healthy participants (50 Malayalam speakers and 50 Kannada speakers) and 20 individuals with mild-to-moderate dysphonia. Voice samples consisted of sustained vowel phonation and connected speech, which were analysed acoustically using AVQI v02.02. Auditory perceptual evaluation was carried out using the Grade parameter of the GRBAS scale. The dysphonic group included individuals with various laryngeal pathologies such as vocal nodules, polyps, cysts, sulcus vocalis, vocal fold paralysis, laryngitis, glottic insufficiency, and muscle tension dysphonia. The results showed significantly higher AVQI values in dysphonic voices compared to healthy voices, and authors reported a normative AVQI value of 3.03 (± 0.32) for combined Malayalam and Kannada speaking adult population with no significant effect of language on AVQI scores.

Vishali (2019) established normative reference data for AVQI v02.02 in Tamil-speaking adults aged between 20 and 50 years. The study was carried out on 136 Tamil-speaking adults, comprising 121 normophonic participants and 15 participants with dysphonia. Acoustic analysis of the recorded voice samples was performed using PRAAT software (version 6.0.28) installed on a Lenovo laptop. The study revealed a normative AVQI value of **2.76** for adult Tamil speakers. This value varied from previously reported Indian and Western thresholds, supporting the language-specific nature of AVQI reference values.

A study by Pebbili et al. (2021) focused on the diagnostic accuracy of AVQI v02.03 in Kannada Speaking adults. The study included 90 voice

samples, comprising 71 dysphonic and 19 normophonic voices. The dysphonic group consisted of individuals between 12 and 82 years of age, whereas the normophonic group had a mean age of 24 years. Acoustic analysis was carried out using Praat software (version 6.0.40) with the AVQI version 02.03 script. Sustained vowel phonation and connected speech samples were analysed, and perceptual severity was rated using the ‘Grade’ parameter of the GRBAS scale. The authors reported an AVQI cut-off value of 2.72 for distinguishing between normal and dysphonic voices with a sensitivity of 52% and a specificity of 73% and AVQI scores increasing with perceptual severity.

Table 2.1

Summary of AVQI Validation and Normative Studies Across Different Languages and Populations

Author & Year	AVQI Version	Language/ Population	Study Type	AVQI Value
Maryn et al. (2010)	v02.02	Dutch	Validation	2.95
Barsties & Maryn (2015)	v02.02	Dutch	Validation	2.80
Barsties & Maryn (2012)	v02.02	German	Validation	2.70
Maryn et al. (2014)	v02.02	German	Validation	3.05
Maryn et al. (2014)	v02.02	French	Validation	3.07
Reynolds et al. (2012)	v02.02	Australian English	Validation	3.46
Maryn et al. (2014)	v02.02	English	Validation	3.25

Hosokawa et al. (2017)	v02.02	Japanese	Validation	3.15
Kim et al. (2018b)	v02.02	Korean	Validation	3.33
Uloza et al. (2018)	v02.02	Lithuanian	Validation	2.97
Kankare et al. (2020)	v02.02	Finnish	Validation	3.09
Zainae et al. (2022)	v02.06	Persian	Validation	3.47–4.04
Pebbili et al. (2021)	v02.03	Kannada	Validation	2.72
Barsties & Maryn (2015)	v03.01	Dutch	Validation	2.43
Barsties & Maryn (2016)	v03.01	Dutch	Validation	2.43
Hosokawa et al. (2019)	v03.01	Japanese	Validation	1.41
Delgado Hernández et al. (2018)	v03.01	Spanish	Validation	2.28
Pommée et al. (2020)	v03.01	French	Validation	2.33
Latoszek et al. (2020)	v03.01	German	Validation	1.85
Englert et al. (2021)	v03.01	Brazilian Portuguese	Validation	1.33
Kim et al. (2020)	v03.01	Korean	Validation	3.15
Fantini et al. (2023)	v03.01	Italian	Validation	2.35
Zainae et al. (2022)	v03.01	Persian	Validation	3.03–3.07
Vishali (2019)	v02.02	Tamil	Normative	Adults: 2.76
Benoy (2017)	v02.02	Kannada & Malayalam	Normative	Adults: 3.03

Jayakumar et al. (2022)	v02.02	Kannada, Malayalam & Tamil	Normative	Adults: 3.03; Pediatrics: 4.02; Older Adults: 3.89
Varna (2024)	v03.01	Kannada & Malayalam	Normative	Pediatrics: 4.02*

* Studies in pediatric population

Note: Validation study, refer to studies that evaluated the diagnostic performance of AVQI in differentiating between normophonic and dysphonic voices, whereas normative studies established reference AVQI values in healthy populations.

2.6 Effect of age and gender on AVQI

Barsties et al. (2017) studied the effects of age and gender on AVQI (v02.02) and DSI. The study comprised 123 normophonic adults between the ages of 20 and 79 years, representing both genders. Although there has been prior research on DSI in this aspect, this was the first study to assess the effects of age and gender on AVQI. The study findings showed that neither of the multiparametric indices showed any significant differences between genders. Similar results were found in Indian studies by Benoy (2017) and Jayakumar et al. (2022), where AVQI values did not differ significantly between genders despite gender-related differences in individual acoustic parameters.

In contrast, age has a clear influence on AVQI thresholds. Adult AVQI v02.02 cut-off values in Indian languages typically range between 2.72 and 3.03 (Benoy, 2017; Vishali, 2019; Pebbili et al., 2021). The Pediatric population has been found to have higher thresholds, with values ranging from 3.74 to 4.02 for AVQI v02.02 (Seshasri, 2018; Jayakumar et al., 2022; Varna,

2024). Older adults showed AVQI values higher than the younger adults with thresholds reported around 3.89 (Jayakumar et al., 2022).

There have been numerous studies conducted on AVQI among different languages in both Western and Indian languages, such as Kannada, Malayalam, and Tamil. However, currently there are no similar studies in the Telugu-speaking population. This means that the reference values of AVQI are not yet clinically available for Telugu-speaking adults.

Therefore, the present study aimed to address this issue through the establishment of normative reference values of AVQI v02.02 among healthy Telugu-speaking adults. The development of population specific normative values is crucial because it allows clinicians to perform accurate clinical evaluations, facilitates meaningful interpretation of AVQI results, and expands the use of AVQI in multilingual clinical settings.

CHAPTER III

METHOD

3.1. Research design

The present study employed a normative research design to investigate the aims and objectives.

3.2. Participants

A total of one hundred and thirty six native Telugu-speaking adults were considered as participants in the present study. The participants were divided into four major groups. Group I, Group II and Group III together consisted of one hundred and twenty three normo-phonic adults recruited randomly from community settings in Andhra Pradesh. The first group (group I) included 39 normo-phonic young adults (17 males and 22 females) aged 20–35 years. The second group (group II) included 43 normo-phonic middle-aged adults (18 males and 25 females) aged 36–50 years. The third group (group III) included 41 normo-phonic older adults (20 males and 21 females) aged 51–65 years. Group IV included 13 adults with dysphonia. Group I, II and III were referred to as control group in the present study, and Group IV as clinical group.

3.2.1. Inclusion criteria

Participants were considered for the study if they met the following inclusion criteria:

- i. Native speakers of Telugu.
- ii. Had learned Telugu as their first language till 10th standard.
- iii. Had used Telugu for at least 15 years.
- iv. Had resided in a Telugu-speaking region for at least 15 years.
- v. Had perceptually normal voice quality as rated as Zero ('0') on overall grade (G) of the GRBAS scale (Hirano et al., 1981) by researcher

3.2.2 Exclusion Criteria

Participants who met the following exclusion criteria, were excluded from the study;

- i. Had a current or recent past history of upper vocal-tract infections, chronic obstructive pulmonary disease, asthma, or other lung infections.
- ii. Had a history of speech, language, or hearing complaints, communication disorders, or neurological or cognitive impairments.
- iii. Had current or past history of alcohol, smoking, or tobacco habits.

3.3. Stimuli

A phonation sample of vowel /a:/ was recorded for a minimum of 5 to 6 seconds, along with a continuous speech/reading sample in Telugu. The initial and final 1 second phonation sample were excluded and middle 3 seconds of phonation sample was considered for analysis. The standardized Telugu passage developed by Savithri & Jayaram (2005) was used for the reading task (continuous speech). Participants were asked to read five sentences; the first and the last sentences were excluded, and the middle three sentences, which contained 56 syllables, were considered for further analysis.

3.4. Instrumentation

Voice samples were recorded using Olympus LS-100 digital voice recorders with a sampling rate of 44.1 KHz and 16-bit resolution. The recorder-to-mouth distance was maintained at approximately 10 cm throughout the recording procedure. Acoustic analysis was carried out using Praat software (version 6.4.42, August 2025 release), which is freely available from the official Praat website, on a Lenovo IdeaPad 3 14ITL6 laptop.

3.5. Procedure

The participants were informed about the objectives, procedures, and tasks beforehand, and their assent (Appendix-A) was obtained prior to the recording session. The study followed the ethical guidelines established by the Ethics Committee for bio-behavioral research involving human subjects of the All India Institute of Speech and Hearing (Venkatesan, 2009). The participants were seated comfortably in an upright position in a noise-

free room during the recording session. The microphone was placed 8–10 cm away from the participant's mouth (Patel et al., 2018) to avoid breathing noise.

First, the participants were made to phonate a sustained vowel /a:/ for a minimum of 5 to 6 seconds. Second, the participants were asked to read the standardized reading passage (Appendix-B) continuously in the Telugu language. Both samples were recorded at the participant's comfortable pitch and loudness level. Each participant was made to familiarize with the reading material before the actual recording. Three trials of a sustained vowel phonations and a continuous speech sample were recorded from each of the participants. The best phonation out of the three trials was considered for further analysis.

3.6. Data Analysis

3.6.1. Data Screening

Prior to acoustic analysis, all recorded samples were screened to ensure technical consistency across recordings. Among the one hundred and twenty three normophonic recordings, thirty one samples were recorded using Recorder 1 in stereo mode, while the remaining ninety two samples were recorded using Recorder 2 in mono mode.

As AVQI analysis is typically carried out using a mono audio signal, the stereo-recorded samples obtained from Recorder 1 were converted to mono format in Praat software before analysis. During preliminary analysis, slight differences in AVQI values were observed between the converted samples and the samples originally recorded in mono mode using Recorder 2. These differences were considered to be potentially related to variations in recording format and the stereo-to-mono conversion process. To maintain uniform recording conditions for acoustic analysis, the thirty one samples recorded using Recorder 1 were not included in the final statistical analysis. Accordingly, the final analysis was carried out using ninety two normophonic samples and thirteen dysphonic samples recorded in mono mode with Recorder 2.

3.6.2. Acoustic Analysis

The stable middle portion of the sustained vowel phonation /a:/ was extracted and labeled as “SV” (sustained vowel). Similarly, the continuous speech sample corresponding to the selected sentences from the standardized reading passage was extracted and labeled as “CS” (continuous speech). The edited samples were saved in .wav format for further analysis.

The saved .wav files were analyzed using Praat software along with the AVQI script developed by Maryn and Weenink (2015) for obtaining AVQI version 02.02 (Barsties & Maryn, 2015). The AVQI script was copied into a text file and saved in “.txt” format. The .wav files and the corresponding script file were opened in Praat software using the “Open Praat script” option. Following execution of the algorithm, AVQI values and the constituent acoustic parameters were obtained. The AVQI script automatically extracted the constituent acoustic parameters, while the software algorithm generated the final AVQI value.

AVQI (v02.02) proposed by Maryn and Weenink (2015) was calculated using the following formula:

$$\text{AVQI} = 9.072 - 0.245\text{CPPS} - 0.161\text{HNR} - 0.470\text{SL} + 6.158\text{ShdB} - 0.071\text{Slope} + 0.170\text{Tilt}$$

The following acoustic measures were extracted and tabulated:

- i. Acoustic Voice Quality Index (AVQI)
- ii. Smoothed Cepstral Peak Prominence (CPPS)
- iii. Harmonic-to-Noise Ratio (HNR)
- iv. Shimmer Local (SL)
- v. Shimmer dB (ShdB)
- vi. Slope of Long-Term Average Spectrum (LTAS)

vii. Tilt of the regression line through the Long-Term Average Spectrum (Tilt)

The AVQI constituent parameters and AVQI scores were tabulated across the four groups. Figure 3.1 presents the output of AVQI analysis obtained from Praat for a normophonic voice sample. Figure 3.2 shows the AVQI output for a dysphonic individual.

Figure 3.1

Graphical representation of AVQI results for a normophonic adult

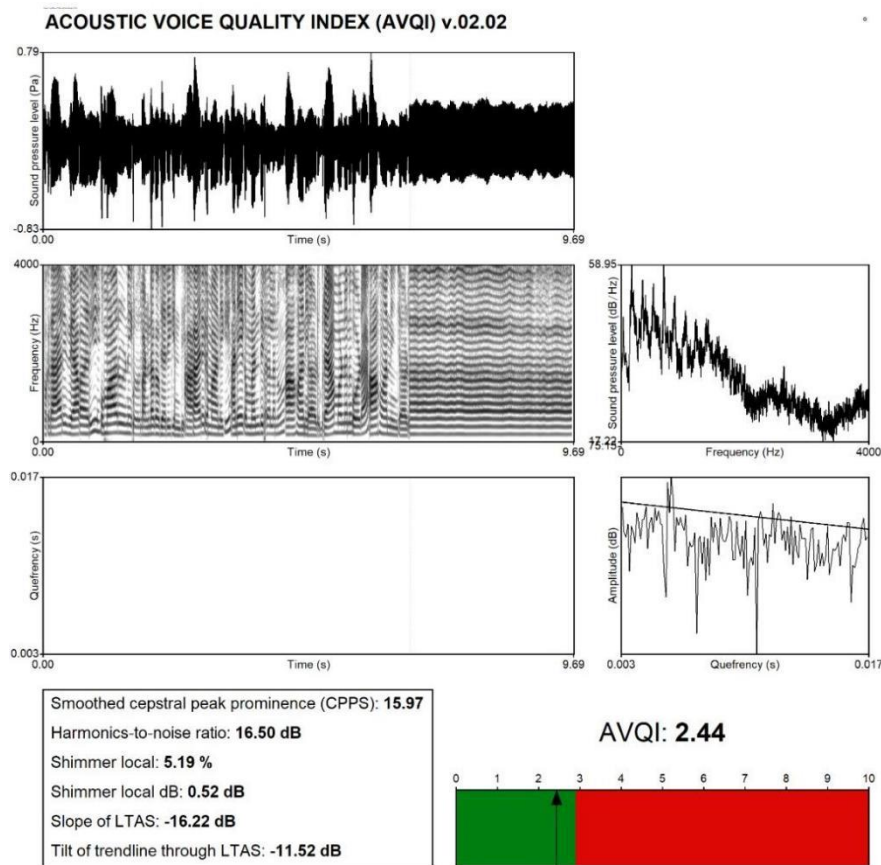
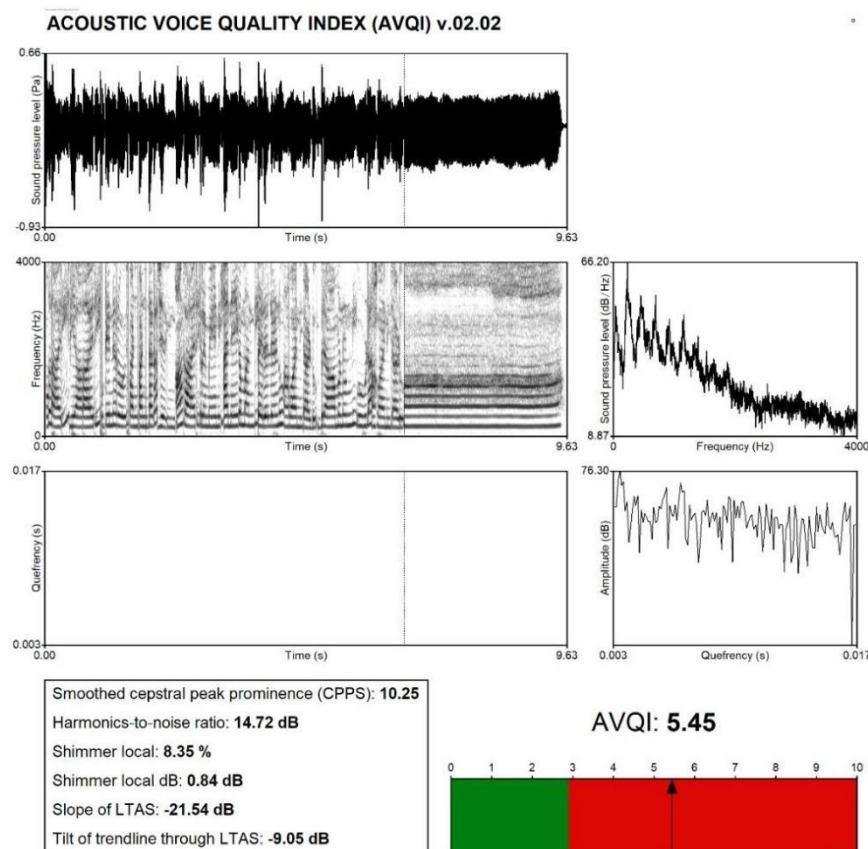


Figure 3.2

Graphical representation of AVQI results for a dysphonic individual.



3.7. Statistical Analysis

The following statistical tests were performed to analyze the data obtained. Shapiro Wilks test of normality was done to check whether the data distributed normally. Descriptive statistics were carried out to document the mean and standard deviation values for AVQI and its constituent parameters for normophonic individuals and individuals with dysphonia. The effect of age and gender on AVQI scores and its constituent parameters were analysed using Two-Way MANOVA test. Post hoc pairwise comparisons with Bonferroni correction were performed to determine which age groups differed significantly from one another on AVQI scores. The difference

between the normophonic group and the dysphonic group on AVQI scores and its constituent parameters was examined using the Mann-Whitney U test. Mann-Whitney U test was done to check the gender differences within the dysphonic group.

CHAPTER IV

RESULTS

The present study investigated the Acoustic Voice Quality Index (AVQI) in native Telugu speakers, with the objectives of establishing normative reference data for AVQI across different age groups, examining the influence of age and gender on AVQI scores, and validating the AVQI in few individuals with dysphonia. A total of ninety-two normophonic individuals participated in the study and were distributed across three age groups: 24 young adults (10 males and 14 females; 20–35 years) as group I, 35 middle-aged adults (12 males and 23 females; 36–50 years) as group II, and 33 older adults (15 males and 18 females; 51–65 years) as group III. In addition, a dysphonic group consisting of 13 individuals also participated in the study as group IV (6 males and 7 females). The AVQI and its constituent parameters, namely CPPS, HNR, Shimmer Local (%), Shimmer Local (dB), Slope of LTAS, and Tilt through LTAS, were extracted from recorded sustained vowel phonation and speech samples using PRAAT software. Age group, gender and phonatory status served as the independent variables, while AVQI and its constituent parameters were treated as the dependent variables. All statistical analyses were performed using IBM SPSS software version 26. Descriptive statistics including mean (M) and standard deviation (SD) were calculated for all measures across groups. The Shapiro-Wilk test of normality was administered on the data across age and gender, where the results indicated a normal distribution of data across all subgroups ($p > 0.05$), except for Shimmer Local (%) among males in the 20–35 years age group ($p = 0.012$). Therefore, Two-way Multivariate Analysis of Variance (MANOVA) with Bonferroni-corrected post hoc comparisons was performed to examine the effects of age and gender on AVQI scores and its constituent parameters. The Mann-Whitney U test was employed to examine the difference between the normophonic and dysphonic groups on AVQI and its constituent

parameters, owing to the unequal sample sizes between the two groups. The details of these analyses are presented below.

4.1.1 Comparison of AVQI across the age groups

Descriptive statistics for AVQI were obtained from ninety two normophonic Telugu-speaking individuals. Table 4.1 presents the mean and standard deviation values of AVQI across the three age groups for both genders.

Table 4.1

Mean (M), standard deviation (SD), Median (MD), Interquartile Range (IQR) values of AVQI in normo-phonic individuals

Groups	Gender	M	SD	MD	IQR
Group I	Males	3.14	0.58	3.39	1.08
	Females	3.22	0.51	3.41	0.98
	Average	3.19	0.53	3.39	0.99
Group II	Males	3.37	0.57	3.38	1.02
	Females	2.89	0.57	2.84	0.74
	Average	3.06	0.60	2.95	0.89
Group III	Males	3.61	0.58	3.69	0.64
	Females	2.92	0.46	2.90	0.92
	Average	3.24	0.61	3.23	0.93

Note: M, Mean; SD, Standard Deviation; Median; IQR, Inter Quartile Range

Across the three age groups, the average mean AVQI was 3.19 (SD = 0.51) for group I, 3.06 (SD = 0.60) for group II, and 3.24 (SD = 0.61) for group III. In group I, females demonstrated slightly higher mean AVQI values. In contrast, males demonstrated higher mean AVQI values than females in group II and group III.

4.2 Comparison of constituent parameters of AVQI across the Age groups

The constituent parameters of AVQI were also examined across three age groups using descriptive statistical measures. The mean (M) and standard deviation (SD) values for parameters including CPPS, HNR, Shimmer Local (%), Shimmer dB, LTAS, and LTAS Tilt obtained from normophonic participants in group I, group II, and group III are summarized in Table 4.2.

Table 4.2 shows the descriptive statistics of the constituent parameters of AVQI across three age groups. CPPS, Shimmer dB, LTAS, and LTAS Tilt showed similar values across groups I, II, and III. In contrast, HNR values were lower in group III, whereas Shimmer Local values were relatively higher in group II and group III compared to group I.

Additionally, a Two-way MANOVA was conducted to evaluate the influence of age on AVQI and its constituent parameters. The findings demonstrated no statistically significant effect of age on AVQI and its constituent parameters ($p > 0.05$). The results of two-way MANOVA for comparison of age as a factor on AVQI and its constituent parameters are represented in table 4.3.

Table 4.2

Mean (M), standard deviation (SD), Median, and interquartile range values of the constituent parameters of AVQI among normophonic individuals across group I, group II, and group III

Groups	Gender	CPPS		HNR		Shimmer Local		Shimmer dB		LTAS		Tilt of LTAS	
		M	SD	M	SD	M	SD	M	SD	M	SD	M	SD
I	Males	14.82	1.11	17.18	1.37	6.45	0.86	0.63	0.07	-21.67	2.23	-11.52	0.56
	Females	13.95	1.07	18.23	1.70	6.35	1.10	0.64	0.08	-20.10	3.23	-11.32	0.67
	Average	14.31	1.15	17.79	1.63	6.39	0.99	0.64	0.07	-20.75	2.91	-11.40	0.62
II	Males	14.60	1.44	16.80	1.52	6.66	1.12	0.64	0.10	-22.3	2.39	-11.07	0.69
	Females	14.86	1.20	17.81	1.73	6.40	1.47	0.64	0.10	-19.89	2.74	-11.77	0.80
	Average	14.77	1.27	17.46	1.71	6.49	1.35	0.64	0.10	-20.73	2.85	-11.53	0.83
III	Males	14.33	1.14	16.05	1.93	7.42	1.70	0.71	0.13	-22.02	2.13	-10.7	0.59
	Females	15.28	0.90	17.88	1.70	6.07	1.01	0.63	0.08	-19.38	2.70	-11.36	0.64
	Average	14.85	1.11	17.05	2.00	6.68	1.51	0.66	0.11	-20.58	2.76	-11.10	0.68

Note: CPPS, Smoothened Cepstral Peak Prominence; HNR, Harmonic to Noise Ratio; dB, decibel; LTAS, Long-term average spectrum; M, Mean ; SD, Standard deviation

4.3 Comparison of AVQI values between genders within each group

Table 4.1 shows the comparison of mean AVQI values between males and females across the three age groups. In group I, females showed slightly higher mean AVQI values than males. In contrast, males recorded higher mean AVQI values than females in both group II and group III.

Further, a Two-way MANOVA was performed to examine the effect of gender on AVQI and its constituent parameters. The analysis revealed a statistically significant main effect of gender on AVQI ($p = 0.003$), HNR ($p = 0.001$), Shimmer Local (%) ($p = 0.044$), and Slope of LTAS ($p = 0.000$) and Tilt through LTAS ($p = 0.018$). Higher mean values of AVQI, Shimmer Local and Tilt through LTAS were observed in males, whereas females demonstrated higher HNR values and slope of LTAS. No significant main effect of gender was observed for CPPS ($p = 0.647$) and Shimmer Local (dB) ($p = 0.220$).

4.4 Effect of age group and gender on AVQI and its constituent parameters

A Two-way MANOVA was performed to examine the effects of age group and gender on AVQI and its constituent parameters.

4.4.1 Effect of age on AVQI and its constituent parameters

The analysis of between-subjects effects indicated no statistically significant effect of age on AVQI and its constituent parameters ($p > 0.05$). The F-values, p-values, and partial eta squared values (results of two-way MANOVA for age comparison) for each parameter are presented in Table 4.3.

Table 4.3

Results of Two-way MANOVA for age comparison on AVQI and its constituent parameters

Parameters	F-value	p-value	Partial Eta Squared
AVQI	0.526	0.59	0.01
CPPS	0.983	0.37	0.02
HNR	1.274	0.28	0.02
Shimmer Local (%)	0.518	0.59	0.01
Shimmer Local (dB)	0.831	0.43	0.01
Slope of LTAS	0.202	0.81	0.00
Tilt through LTAS	2.648	0.07	0.05

4.4.2 Effect of gender on the constituent parameters of AVQI

The analysis revealed a statistically significant effect of gender on AVQI ($p = 0.003$), HNR ($p = 0.001$), Shimmer Local (%) ($p = 0.044$), Slope of LTAS ($p = 0.000$), and Tilt through LTAS ($p = 0.018$). Higher mean values of AVQI, Shimmer Local (%), and Tilt through LTAS were observed in males, whereas females demonstrated significant higher HNR values and Slope of LTAS. No statistically significant effect of gender was observed for CPPS and Shimmer Local (dB) ($p > 0.05$). The F-values, p-values, and partial eta squared values are summarized in Table 4.4.

Table 4.4

Results of Two-way MANOVA for gender comparison on AVQI and its constituent parameters

Parameter	F-value	p-value	Partial Eta Squared
AVQI	9.244	0.00*	0.097
CPPS	0.211	0.64	0.002
HNR	12.458	0.00*	0.127
Shimmer Local (%)	4.178	0.04*	0.046
Shimmer Local (dB)	1.525	0.22	0.017
Slope of LTAS	15.164	0.00*	0.150
Tilt through LTAS	5.860	0.01*	0.064

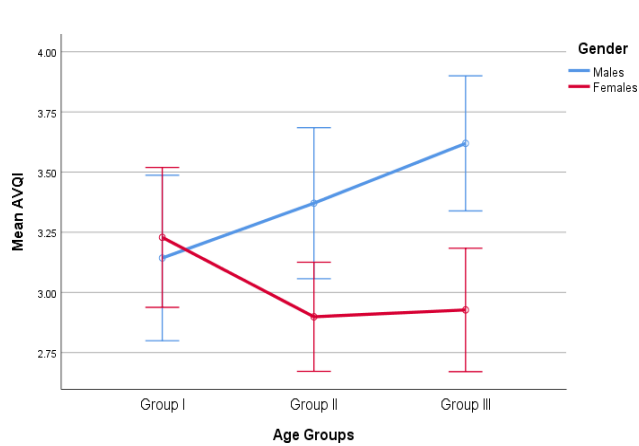
* indicates significant difference at 0.05 level

4.4.3 Interaction effect of age and gender

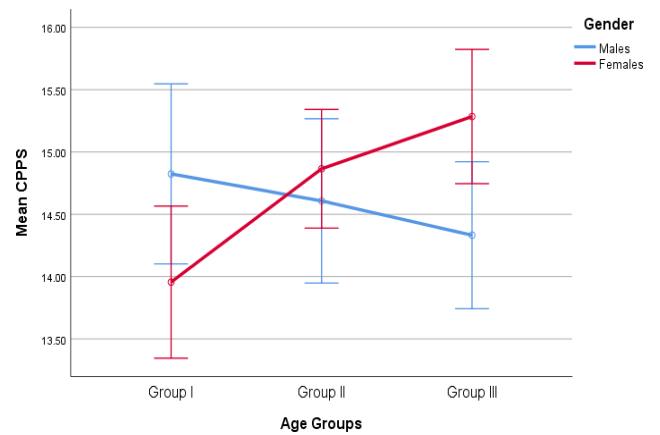
A statistically significant interaction effect between age and gender was observed for AVQI, CPPS, and LTAS Tilt ($p < 0.05$). Specifically, a significant difference was observed for the age*gender interaction effect for AVQI, $F(2) = 3.531$ ($p < 0.05$); CPPS, $F(2) = 4.284$ ($p < 0.05$); and LTAS Tilt, $F(2) = 3.259$ ($p < 0.05$). However, no significant interaction effect was observed between age and gender for HNR, $F(2) = 0.573$ ($p > 0.05$); Shimmer Local (%), $F(2) = 2.099$ ($p > 0.05$); Shimmer Local (dB), $F(2) = 1.871$, ($p > 0.05$); and Slope of LTAS, $F(2) = 0.311$, ($p > 0.05$).

Figure 4.1

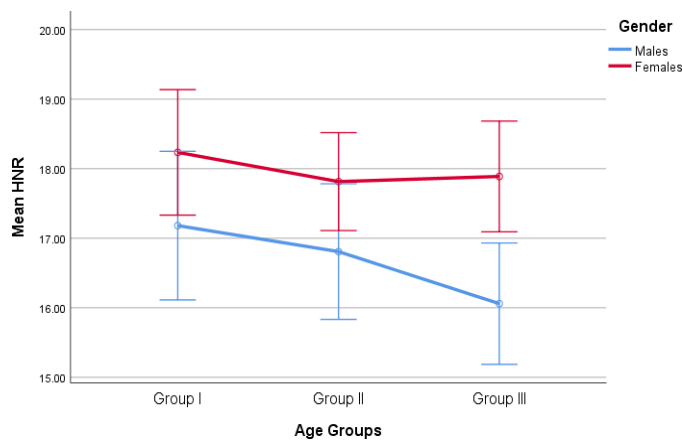
Mean AVQI, CPPS, HNR and shimmer Local (%) values between Males and Females across three groups



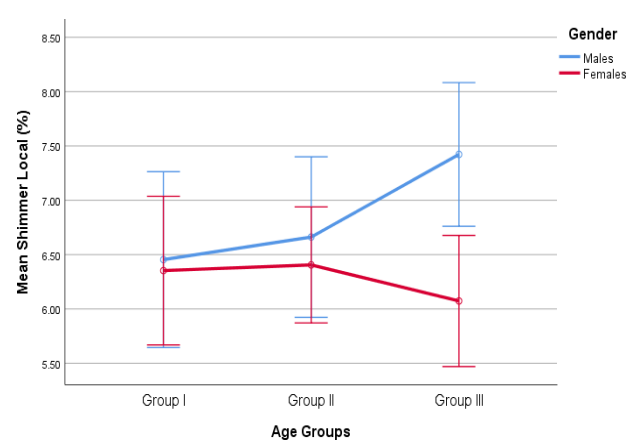
A) Mean AVQI between males and females across three groups



B) Mean CPPS between males and females across three groups



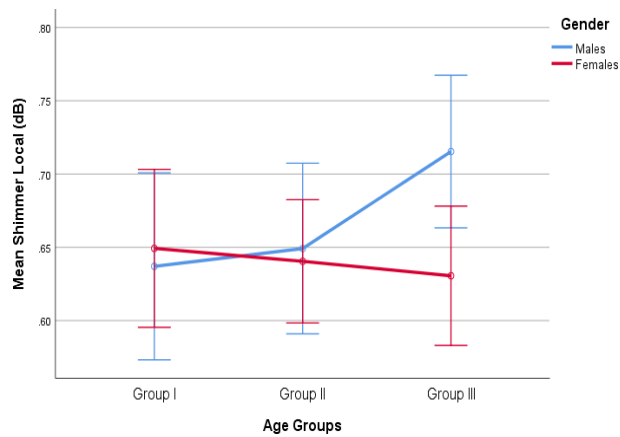
C) Mean HNR between males and females across three groups



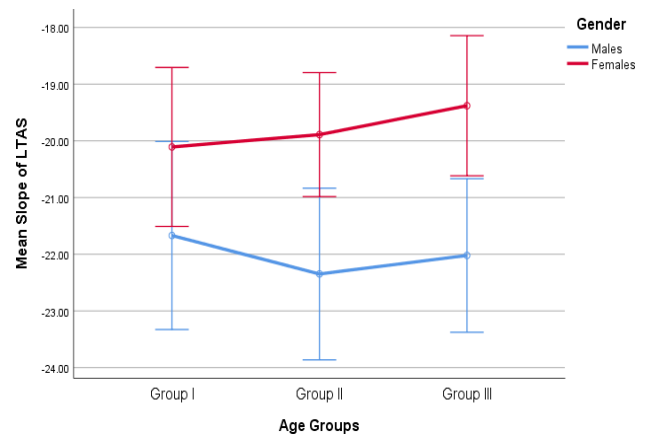
D) Mean Shimmer Local (%) between males and females across three groups

Figure 4.2

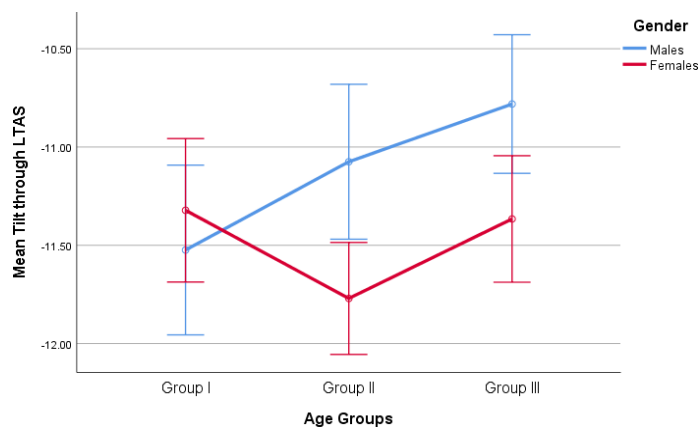
Mean shimmer Local dB ,slope of LTAS, and Tilt through LTAS values between males and females across three groups.



A) Mean Shimmer Local (dB) between males and females across three groups



B) Mean Slope of LTAS between males and females across three groups



C) Mean Tilt through LTAS between males and females across three groups

[illegible]

4.6 Comparison of AVQI and its constituent parameters between normophonic and dysphonic Individuals

The comparison of AVQI and its constituent acoustic parameters was extracted and tabulated for thirteen dysphonic individuals (group IV). The descriptive statistics, including mean (M) and standard deviation (SD), obtained for the dysphonic group (group IV) are presented in Table 4.6.

Table 4.6

Mean (M) and Standard Deviation (SD) values of AVQI and its constituent parameters among dysphonic individuals (group IV)

Paramters	M	SD
AVQI	5.33	1.73
CPPS	10.26	3.33
HNR	13.47	4.37
Shimmer Local (%)	10.81	3.86
Shimmer Local (dB)	1.02	0.29
Slope of LTAS	-20.96	3.83
Tilt through LTAS	-10.34	1.43

Note: M, Mean; SD, Standard Deviation; AVQI, Acoustic Voice Quality Index; CPPS, Smoothened Cepstral Peak Prominence; HNR, Harmonic to Noise Ratio; dB, decibel; LTAS, Long-term average spectrum

The obtained values showed that the dysphonic group demonstrated higher mean scores for AVQI, Shimmer Local (%), Shimmer Local (dB) and Tilt through LTAS when compared with the normophonic group. In contrast, lower mean values were observed for CPPS and HNR in the dysphonic participants (group IV). The Slope of LTAS parameter exhibited nearly equivalent values in the two groups. Graphical representation of comparison of Mean values of AVQI and its constituent acoustic parameters between normophonic and dysphonic individuals is shown in Figure 4.3.

Figure 4.3 A.

Comparison of Mean values of AVQI, CPPS, HNR, Shimmer (%), and Shimmer (dB) among normophonic males, normophonic females, and dysphonic individuals.

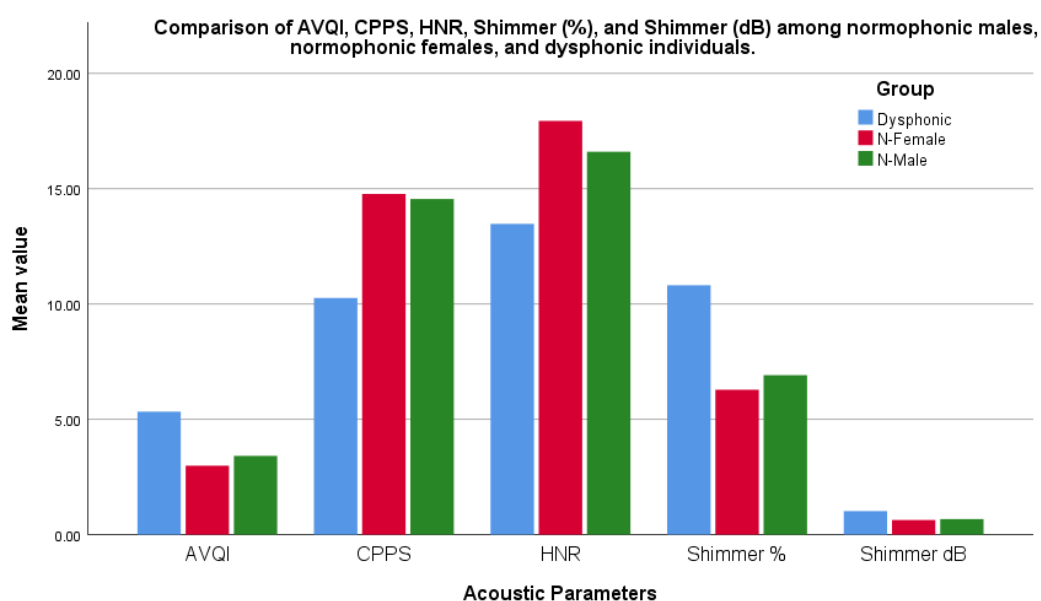
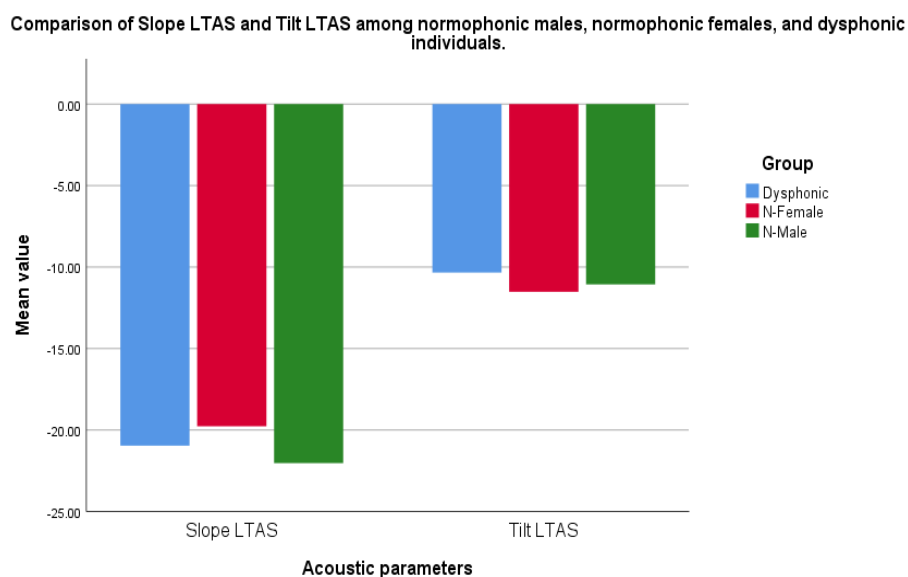


Figure 4.3.B.

Comparison of Mean values of slope of LTAS and Tilt through LTAS among normophonic males, normophonic females, and dysphonic individuals



To examine group differences statistically, the Mann–Whitney U test was applied for AVQI and its constituent parameters. The analysis demonstrated statistically significant differences between the normophonic and dysphonic groups for AVQI, CPPS, HNR, Shimmer Local (%), Shimmer Local (dB), and Tilt through LTAS ($p < 0.05$). However, no statistically significant difference was identified for Slope of LTAS ($p = 0.782$). The findings of the Mann–Whitney U test are presented in Table 4.7.

Table 4.7

Results of Mann–Whitney U test comparing normophonic and dysphonic groups on AVQI and its constituent parameters

Acoustic Measures	Mann–Whitney U	Z	p-value
AVQI	74.00	5.09	< 0.00*
CPPS	50.50	5.32	< 0.00*
HNR	218.00	3.69	< 0.00*
Shimmer Local (%)	108.00	4.76	< 0.00*
Shimmer Local (dB)	74.500	5.09	< 0.00*
Slope of LTAS	569.500	0.27	0.782
Tilt through LTAS	295.500	2.94	0.003*

‘*’ indicates significant difference at 0.05 level

4.7 Establishing normative data for AVQI and its constituent parameters in Telugu speaking adults

Table 4.8 presents the normative reference data of AVQI and its constituent parameters for native Telugu speaking normophonic adults across three age groups: young adults (group I), middle-aged adults (group II), and older adults (group III) separately for males and females. Since the AVQI scores did not show a significant effect of age, the reference measures are presented collectively across all three age groups in Table 4.8 . However, as gender was found to have a significant influence on AVQI scores, separate reference values are provided for males and females. These values may serve as preliminary normative reference data for the assessment of voice quality in native Telugu speaking adults using the Acoustic Voice Quality Index.

The overall mean AVQI for native Telugu-speaking normophonic adults aged 20–65 years was 3.41 (SD = 0.59) for males and 2.99 (SD = 0.53) for females, across all three age groups combined. The overall median AVQI for native Telugu-speaking normophonic adults aged 20–65 years was 3.49 (IQR= 0.91) for males and 3.05 (IQR = 0.88) for females, across all three age groups combined.

Table 4.8

Reference scores of AVQI (M and SD) and its constituent parameters in native Telugu speaking normophonic adults across three age groups and between males and females

Parameters	20–35 years (Group I)		36–50 years (Group II)		51–65 years (Group III)		Average	
	Males	Females	Males	Females	Males	Females	Males	Females
	M (SD)	M (SD)	M (SD)	M (SD)	M (SD)	M (SD)	M (SD)	M (SD)
AVQI	3.14 (0.58)	3.23 (0.52)	3.37 (0.58)	2.90 (0.57)	3.62 (0.58)	2.93 (0.46)	3.41 (0.59)	2.99 (0.53)
CPPS	14.82 (1.12)	13.96 (1.08)	14.61 (1.45)	14.87 (1.20)	14.33 (1.14)	15.28 (0.91)	14.55 (1.22)	14.77 (1.17)
HNR	17.18 (1.38)	18.23 (1.71)	16.81 (1.52)	17.81 (1.74)	16.06 (1.93)	17.89 (1.70)	16.60 (1.69)	17.94 (1.69)
Shimmer Local (%)	6.45 (0.87)	6.35 (1.10)	6.66 (1.13)	6.41 (1.47)	7.42 (1.71)	6.07 (1.01)	6.91 (1.37)	6.28 (1.23)
Shimmer Local (dB)	0.64 (0.07)	0.65 (0.08)	0.65 (0.11)	0.64 (0.11)	0.72 (0.13)	0.63 (0.09)	0.67 (0.11)	0.63 (0.09)
Slope of LTAS	21.67 (2.23)	-20.11 (3.24)	-22.35 (2.39)	-19.89 (2.75)	-22.02 (2.14)	-19.38 (2.70)	-22.03 (2.20)	-19.77 (2.82)
Tilt through LTAS	-11.52 (0.57)	-11.32 (0.67)	-11.08 (0.70)	-11.77 (0.81)	-10.78 (0.60)	-11.37 (0.65)	-11.07 (0.67)	-11.52 (0.74)

Note: M, Mean; SD, Standard deviation (depicted in parenthesis)

Table 4.9

Median and Interquartile Range scores of AVQI and its constituent parameters in native Telugu speaking normophonic adults across three age groups and between males and females

Parameters	20–35 years (Group I)		36–50 years (Group II)		51–65 years (Group III)		Average	
	Males	Females	Males	Females	Males	Females	Males	Females
AVQI	3.39 (1.08)	3.41 (0.98)	3.38 (1.02)	2.84 (0.74)	3.69 (0.64)	2.90 (0.92)	3.49 (0.91)	3.05 (0.88)
CPPS	14.71 (1.59)	13.66 (1.95)	14.43 (2.42)	14.97 (1.63)	14.43 (1.42)	15.02 (1.37)	14.53 (1.81)	14.55 (1.65)
HNR	17.50 (2.67)	18.37 (2.84)	16.98 (2.63)	18.02 (2.30)	16.76 (3.09)	17.78 (1.94)	17.08 (2.80)	18.06 (2.36)
Shimmer Local (%)	6.26 (0.98)	6.19 (1.55)	6.47 (1.08)	6.51 (2.30)	7.42 (1.56)	5.82 (1.51)	6.72 (1.21)	6.17 (1.79)
Shimmer Local (dB)	0.61 (0.13)	0.64 (0.11)	0.64 (0.12)	0.63 (0.13)	0.71 (0.10)	0.60 (0.12)	0.65 (0.12)	0.63 (0.12)
Slope of LTAS	-21.92 (3.54)	-19.61 (4.16)	-22.64 (4.70)	-19.91 (4.47)	-22.44 (3.41)	-18.56 (4.45)	-22.34 (3.88)	-9.36 (4.36)
Tilt through LTAS	11.55 (0.99)	-11.30 (1.25)	-11.29 (0.84)	-11.66 (1.07)	-10.86 (0.71)	-11.43 (0.73)	-11.24 (0.85)	-11.47 (1.02)

Note: Values in parenthesis indicates Inter quartile range

CHAPTER V

DISCUSSION

The present study investigated the Acoustic Voice Quality Index (AVQI v02.02) in native Telugu-speaking adults with the following objectives: to establish normative reference data for AVQI in Telugu-speaking adults aged 20-65 years, to examine the effect of age and gender on AVQI and its constituent parameters, and to validate the obtained AVQI value in a few dysphonic individuals. The following sections discuss the above objectives based on the findings of the study.

5.1. AVQI in Telugu-speaking adults

The overall mean AVQI value obtained for native Telugu-speaking adults aged 20–65 years was $M = 3.41$ ($SD = 0.59$) for males and $M = 2.99$ ($SD = 0.53$) for females. The AVQI values obtained in the present study were similar to those reported in other Indian language populations, including Malayalam-, Kannada-, and Tamil-speaking adults (Benoy & Jayakumar, 2017; Vishali, 2019). These findings suggest that AVQI values are generally comparable across Dravidian languages. The current findings are also in agreement with those reported in several non-Indian language populations, including Dutch with an AVQI value of 2.80 (Barsties & Maryn, 2015), German with 2.70 (Barsties & Maryn, 2012), French with 3.07 (Maryn et al., 2014), English with 3.25 (Maryn et al., 2014), Australian English with 3.46 (Reynolds et al., 2012), Japanese with 3.15 (Hosokawa et al., 2017), Korean with 3.33 (Kim et al., 2018b), Lithuanian with 2.97 (Uloza et al., 2018), Finnish with 3.09 (Kankare et al., 2019), and Persian with 3.47 (Zainae et al., 2022). Slight variation in normative AVQI values can be observed across language populations. These differences may be influenced by linguistic factors such as

phonological and prosodic characteristics, as well as methodological factors including sample size, participant age range, recording procedures, speech materials, equipment employed and analysis methods. Maryn et al. (2014) reported no significant cross-linguistic differences across English, Dutch, French, and German populations. Overall, the Telugu AVQI values obtained in the present study fall within the range reported in both Indian and non-Indian language populations.

5.2. Effect of age on AVQI and its constituent parameters

The findings of the present study revealed no significant effect of age on AVQI among native Telugu-speaking adults aged 20–65 years. That is, young adults aged between 20-35 years, then middle-aged adults aged between 36-50 years, and then older adults aged between 51 to 65 years did not differ significantly on AVQI and its constituent parameters (except tilt through LTAS). This suggests that AVQI remained relatively stable across the adult age range examined in the study. The current findings are in agreement with earlier research noting little to no influence of age on AVQI. Latoszek et al. (2019) found that AVQI was not significantly associated with age in normophonic adults aged 20–79 years. Similar results have been observed among some Indian populations. Benoy and Jayakumar (2017) found no significant age-related differences in AVQI among Malayalam- and Kannada-speaking adults aged between 20-50 years of age and Vishali (2019) reported similar results for Tamil-speaking adults aged between 20-50 years of age. Jayakumar et al. (2022) noted that AVQI remained relatively stable within the adult age range between 20-50 years of age. However, significant age-related effects were observed only when participants across the entire lifespan(wide spectrum of age range, 8-70 years of age), including paediatric and geriatric groups, were considered. The systematic review by Batthyany et al. (2022) further concluded that age has little or no significant influence on AVQI.

Although ageing is associated with well documented structural and functional changes in the larynx such as changes in the vocal fold tissue, muscle function, and vibratory behaviour, these changes do not appear to produce measurable differences in AVQI among adults aged 20 to 65 years. Yamauchi et al. (2015) demonstrated that older adults show greater vocal fold vibratory asymmetry than younger adults, and therefore, some age-related changes in physiology will definitely show changes at a biomechanical level. However, Santos et al. (2023) found no significant age-related differences in shimmer, jitter, HNR, or several other acoustic parameters across adults aged 30 to 79 years, and Rojas et al. (2020) concluded that meaningful perturbation changes generally become more apparent only after 70 years of age. The current results support this pattern, indicating that age-related physiological changes between the ages of 20 and 65 years might not be sufficiently enough to result in significant differences or changes in AVQI or most of its constituent parameters.

The constituent acoustic measures showed a similar pattern as AVQI. CPPS, Shimmer Local (dB) and Slope of LTAS values were mostly similar across the three age groups. The older age groups had higher Shimmer Local (%) values and lower HNR values, but these differences were not statistically significant. Similar findings have been reported by Benoy and Jayakumar (2017) and Santos et al. (2023), who found no significant age-related differences in shimmer %, shimmer (dB) and HNR across adults aged 30–79 years. The relatively stable CPPS values across age groups suggest that the periodicity and harmonicity of the voice were largely maintained from 20 years to 65 years of age. Likewise, the non-significant differences observed for HNR and shimmer measures indicate that age-related changes in noise components and amplitude perturbation were minimal within the age range studied. The stable slope of LTAS further suggests that the overall distribution of spectral energy across different frequencies in the voice remained largely unchanged from

adulthood to older adults (up to 65 years of age). These results imply that, although some acoustic measures may exhibit slight age-related trends, they are unlikely to be significant enough to result in differences in the study population.

A notable variation was observed for the Tilt through LTAS parameter, which showed a significant difference between group II (36–50 years) and group III (51–65 years) in the post hoc analysis. This finding can possibly be the result of subtle age-related changes in the vocal folds. Rojas et al. (2020) reported that ageing is linked with reductions in connective tissue components and hyaluronic acid, which can impact the biomechanical characteristics of the vocal folds. Such changes may be reflected more easily in spectral measures like LTAS than in perturbation-based measures. However, despite this isolated difference in Tilt through LTAS, no significant effect of age was observed for overall AVQI and other parameters, suggesting that age-related changes within the adult age range were not sufficient to substantially influence the composite Acoustic voice quality index.

5.3. Effect of gender on AVQI and its constituent parameters

A significant effect of gender was observed on AVQI in the present study, with males showing higher AVQI values than females. This finding differs from several earlier studies that reported no significant gender effect on AVQI (Latoszek et al., 2019; Batthyany et al., 2022; Benoy & Jayakumar, 2017; Hippargekar et al., 2022; Shabnam & Pushpavathi, 2022; Vishali, 2019). However, significant gender differences were also observed in several constituent parameters of AVQI in the present study, including HNR, Shimmer Local (%), Slope of LTAS, and Tilt through LTAS. Since AVQI is a composite measure derived from multiple acoustic parameters, the cumulative influence of these gender-related differences may have contributed to the significant difference observed in the overall AVQI score.

Jayakumar et al. (2022) similarly noted that individual acoustic parameters may demonstrate gender sensitivity even when the composite AVQI does not, and that the overall AVQI outcome depends on the combined contribution of its constituent measures.

Among the constituent parameters, a significant gender difference was observed for Shimmer Local (%), with males showing higher values than females. Shimmer Local reflects amplitude perturbation during phonation, and higher values indicate greater variability in vocal fold vibration. The higher shimmer values observed in males may be related to anatomical and physiological differences in voice production. Yamauchi et al. (2015) reported larger vocal fold vibration amplitudes in males, while Zhang (2021) reported greater airflow and vocal fold vibration amplitudes in adult males due to differences in vocal fold dimensions. Similar findings have also been reported in previous studies (Rojas et al., 2020; Santos et al., 2023; Shabnam & Pushpavathi, 2022). Although Shimmer Local (dB) was also higher in males, the difference was not statistically significant, suggesting that this measure may be relatively less sensitive to gender-related differences than Shimmer Local (%).

HNR was significantly higher in females than in males. Harmonic to Noise Ratio is the ratio of harmonic energy to noise in the voice signal; higher HNR values indicate a clearer and more periodic voice. The lower HNR observed in males may be due to differences in vocal fold vibration and glottal closure during phonation, which can lead to increased noise in the voice signal (André et al., 2022). Goy et al. (2013) also reported higher HNR values in females and similar findings have been observed in studies investigating AVQI and its constituent parameters (Benoy & Jayakumar, 2017; Jayakumar et al., 2022). The present findings are therefore consistent with the above mentioned studies showing higher HNR values in females. However, not all

studies have found a significant gender difference in HNR. Hippargekar et al. (2022) reported no gender difference in HNR among Indian normophonic adults, and Teixeira and Fernandes (2014) found gender differences only for jitter, but not for HNR. Different results reported across studies might be attributed to variations in sample size, age range, language background, and acoustic analysis procedures.

In the current study, the slope of LTAS demonstrated a significant gender effect, with males having more negative slope values than females, indicating poorer voice quality in males than in females. The LTAS slope reflects the distribution of spectral energy from lower to higher frequencies, with males having a more negative slope indicating a greater concentration of energy in the lower frequency region. This is in agreement with known acoustic differences between male and female voices, with males having a lower fundamental frequency and longer vocal fold length, resulting in greater energy concentration in the lower frequency range (Zhang, 2021).

CPPS did not differ significantly between males and females in the current study. CPPS determines the strength of the cepstral peak and the degree of periodicity in the voice signal. Benoy and Jayakumar (2017) and Jayakumar et al. (2022) found no significant gender differences, which is consistent with the current findings, although Shabnam and Pushpavathi (2022) reported higher CPPS values in males than females. The conflicting results from various studies indicate that gender may have little effect on CPPS and may differ depending on the population and study samples.

The significant gender differences observed in HNR, Shimmer Local (%), Slope of LTAS, and Tilt through LTAS may have contributed to the gender difference in AVQI. However, as previous studies have generally reported little or no effect of gender on AVQI, further research is needed with larger sample size among Telugu speaking participants to better understand the influence of gender.

5.4. Interaction effect of age and gender on AVQI and its constituent parameters

A significant interaction between age group and gender was observed for a few parameters like AVQI, CPPS, and Tilt through LTAS, suggesting that age-related changes in these measures were not the same for males and females. Although age alone did not have a significant effect on AVQI, and gender effects were observed only for some parameters, the interaction findings indicate that the influence of age may differ between men and women. Previous studies have also suggested that vocal ageing does not occur in the same way across genders (Goy et al., 2013; Rojas et al., 2020). The interaction effects observed for CPPS and Tilt through LTAS further suggest that age-related changes in cepstral and spectral voice measures may vary between males and females. However, studies examining age and gender interactions in AVQI are limited, making it difficult to compare the present findings with earlier research. Therefore, further studies with larger samples are needed to better understand these interaction effects.

5.5. Comparison of AVQI and its constituent parameters between normophonic and dysphonic Individuals

AVQI is a sensitive tool for identifying pathological voice. The current study found significantly higher AVQI scores in dysphonic individuals than in normophonic individuals, indicating poorer overall voice quality in the dysphonic group. This finding is in agreement with previous studies that have reported higher AVQI values in individuals with dysphonia (Benoy & Jayakumar, 2017; Vishali, 2019; Batthyany et al., 2022). The strong discriminative ability of AVQI has also been determined

through the pooled AUC, sensitivity, and specificity values reported by Batthyany et al. (2022) and Jayakumar and Benoy (2022).

CPPS and HNR were significantly lower in the dysphonic group than in the normophonic group. Lower CPPS and HNR values in dysphonic individuals indicate reduced voice periodicity and increased noise in the voice signal, reflecting the irregular vocal fold vibration commonly associated with voice disorders. Similar findings have been reported in previous studies by Heman-Ackah et al. (2014), and André et al. (2022). Heman Ackah et al. (2014) reported that the mean CPPS value was significantly lower in dysphonic voices (2.57 ± 1.05) than in normal voices (4.77 ± 0.97), demonstrating the usefulness of CPPS in differentiating normal and dysphonic voice samples. Likewise, André et al. (2022) observed that pathological voices exhibited lower HNR values than healthy voices.

Shimmer Local (%) and Shimmer Local (dB) were significantly higher in the dysphonic group. Increased shimmer values indicate greater cycle-to-cycle amplitude variation, which is characteristic of dysphonic voices and may be associated with hoarseness, breathiness, and vocal instability (Teixeira et al., 2013). Similar findings have been reported by Hippargekar et al. (2022) who observed significantly elevated shimmer local (%) values in pathological voices (4.86%-6.59%) compared to normal healthy voices (2.52%-3.23%) across the phonation of vowels /a/, /i/, and /u/.

The dysphonic group also demonstrated higher Tilt through LTAS values, indicating changes in the spectral characteristics of the voice associated with vocal pathology. However, Slope of LTAS remained similar across the

two groups. A comparable observation was reported by Benoy and Jayakumar (2017), indicating that this parameter may be less useful than other AVQI components in distinguishing normophonic and dysphonic voices, as it appears to have limited sensitivity to voice quality changes associated with dysphonia.

CHAPTER VI

SUMMARY AND CONCLUSIONS

The human voice is multidimensional, and describing its quality with a single acoustic parameter is not enough. Researchers found that no single acoustic measure could fully reflect how a listener perceives the overall quality of a voice. This resulted in the development of multiparametric measures, which combine multiple acoustic parameters into a single value. The Dysphonia Severity Index was among the first such measures (Wuyts et al., 2000). However, it required parameters such as the highest frequency and the lowest intensity, which are time-consuming to obtain and need high-quality instrumentation, making it difficult to apply in routine clinical practice. Also, the DSI was derived primarily from sustained vowel phonation and did not include connected speech samples in its analysis. To address these limitations, researchers aimed for a measure that was simple and easy to use.

Maryn et al. (2010) introduced the Acoustic Voice Quality Index (AVQI) to address this need. It combines a sustained vowel and a continuous speech sample and gives a single score that reflects the overall quality of the voice. The AVQI is calculated from six acoustic parameters, namely smoothed cepstral peak prominence (CPPS), harmonics-to-noise ratio (HNR), shimmer local (%), shimmer local (dB), the slope of the long-term average spectrum, and its tilt. The AVQI can be computed using the freely available Praat software, has been found to correlate well with perceptual ratings of voice. Because of these advantages, it is widely used. Reference values for AVQI have been

reported in many languages, such as Dutch, German, French, English, Korean, Lithuanian, Finnish, and Persian. In Indian languages, normative values for AVQI have been established for Malayalam and Kannada (Benoy, 2017) and for Tamil (Vishali, 2019). To date, however, no study has reported normative reference data for AVQI version 02.02 in Telugu-speaking adults.

Therefore, the present study was carried out to establish normative reference data for AVQI (version 02.02) in normophonic Telugu-speaking adults between 20 and 65 years of age. The study also investigated how age and gender influenced AVQI and its parameters, and validated the obtained AVQI normative value in a few individuals with dysphonia.

A normative research design was used for the study. A total of one hundred and thirty six native Telugu-speaking adults took part in the study. They were divided into four groups. Groups I, II, and III included one hundred and twenty three normophonic adults who were recruited from community settings in Andhra Pradesh, namely 39 young adults (17 males and 22 females; 20-35 years) as group I, 43 middle-aged adults (18 males and 25 females; 36-50 years) as group II, and 41 older adults (20 males and 21 females; 51-65 years) as group III. Group IV included 13 adults with dysphonia (6 males and 7 females; 20-65 years). To keep the recording format the same for all samples, the samples that are recorded in stereo mode were removed during data screening. The final analysis was therefore carried out on 92 normophonic samples distributed across three age groups: 24 young adults (10 males and 14 females; 20–35 years) as group I, 35 middle-aged adults (12 males and 23 females; 36–50 years) as group II, and 33 older adults (15 males and 18 females; 51–65 years) as group III and 13 dysphonic samples (6 males and 7 females; 20-

65 years) as group IV all recorded in mono mode.

Each participant phonated the vowel /a:/ for 5 to 6 seconds, and the middle 3 seconds were extracted and used for further analysis. The participants also read the standardized Telugu passage (Savithri & Jayaram, 2005), from which the middle three sentences which contained 56 syllables were used as the continuous speech sample. The recordings were made in a quiet room using an Olympus LS-100 recorder (44.1 KHz sampling rate, 16-bit resolution), keeping a distance of about 10 cm between the mouth and the microphone. The samples were analysed in Praat software (version 6.4.42), installed on a Lenovo IdeaPad 3 14ITL6 laptop, using the AVQI version 02.02 script (Maryn & Weenink, 2015). The following measures were obtained:

- Acoustic Voice Quality Index (AVQI)
- Smoothed Cepstral Peak Prominence (CPPS)
- Harmonic-to-Noise Ratio (HNR)
- Shimmer Local (SL)
- Shimmer Local dB (ShdB)
- Slope of the Long-Term Average Spectrum (LTAS)
- Tilt through the Long-Term Average Spectrum (Tilt)

The Shapiro-Wilk test was used to check whether the data were normally distributed, and descriptive statistics (mean and standard deviation) were calculated. A two-way MANOVA with Bonferroni post hoc tests was used to find out whether age and gender influenced the AVQI scores and the related parameters, and Mann-Whitney U test was used to compare the normophonic and dysphonic groups.

According to the findings, normophonic Telugu-speaking adults

between the ages of 20 and 65 years had mean AVQIs of 3.41 (SD = 0.59) for men and 2.99 (SD = 0.53) for women. The AVQI values for Telugu speakers fall within the range reported in previous studies, as these values are comparable to those reported for other Indian and non-Indian languages.

In the present study, neither the AVQI nor its parameters were influenced by age. This shows that AVQI stayed relatively stable throughout the adult age groups till 65 years of age examined in the study. The post hoc test revealed that group II and group III differed only in Tilt through LTAS parameter, which had no impact on the overall AVQI. This suggests that the voice changes that occur between 20 and 65 years of age may not have significant changes that can cause differences in AVQI

Gender, however, showed a significant effect. Males had higher values for AVQI, Shimmer Local (%), and tilt through LTAS, whereas females had higher HNR values and slope of LTAS. Males and females had similar CPPS and Shimmer Local (dB) values. Because AVQI is composed of several parameters, the combined effect of these gender differences could have resulted in the observed difference in the overall AVQI. A significant interaction between age and gender was also discovered for AVQI, CPPS, and Tilt of LTAS but further research involving larger sample size is needed to confirm these trends.

The dysphonic group had a significantly higher mean AVQI (5.33) than the normophonic group, indicating poorer overall voice quality. Additionally, the dysphonic group had significant lower CPPS and HNR values and significant higher shimmer values, suggesting a less periodic, noisier, and more unstable vocal signal. Differences were also observed in the LTAS

parameters, with LTAS slope showing comparable values between the groups(i.e no difference between normals and dysphonics), whereas tilt through LTAS is significantly lower in normophonic compared to dysphonic participants. Collectively, these findings suggest that AVQI and its constituent acoustic measures may be useful in distinguishing between normophonic and dysphonic voices. More dysphonic participants are warranted in future studies to establish the Receiver Operating Characteristic curve and to calculate sensitivity and specificity of AVQI and its related parameters in Telugu speaking population.

The comparison between the normophonic and dysphonic groups involved relatively small number of participants and should therefore be interpreted with caution. Future studies involving larger and more diverse samples would be valuable for further examining these findings. Nevertheless, the present study provides reference values for AVQI (version 02.02) in Telugu-speaking adults and adds to the existing literature on AVQI in the Indian population. The observed gender-related differences in AVQI also suggest the need for further exploration of gender-specific reference values. Overall, the study provides preliminary data on AVQI in Telugu-speaking adults for the age range 20- 65 years, which may be helpful in conducting future research and clinical investigations.

The results of the present study revealed several points of interest;

First, the overall median AVQI value obtained for normophonic Telugu-speaking adults of age range between 20-65 years is 3.49 (IQR: 0.91) for males and 3.05 (IQR: 0.88) for females. The study further, highlights, the mean value of AVQI in males is 3.41 (SD = 0.59) and 2.99 (SD = 0.53) in females. The obtained AVQI value between 20 and 65 years old Telugu

speaking normophonic individuals are in line with the published AVQI values across languages.

Second, the AVQI values showed minimal variation across three age groups (young, middle aged, and older adults), though there was no statistical significant differences between the age groups. The present study did not find age as a factor on AVQI and its constituent parameters in Telugu-speaking adults. However, one of the AVQI constituent parameter, Tilt through LTAS showed a significant difference between middle-aged (group II) and older adults (group III), which needs to be investigated further with large sample size.

Third, gender had a significant influence on AVQI value, where males had higher AVQI value than females. Additionally, gender significantly influenced several constituent parameters of AVQI including HNR, Shimmer Local (%), Slope of LTAS, and Tilt through LTAS. It was found that, HNR and Slope of LTAS were significantly higher in females and Shimmer Local (%), and Tilt through LTAS were significantly higher in males. On the other hand, CPPS and Shimmer Local (dB) were found to be not influenced by gender. However, such gender differences in AVQI and its constituent parameters need further exploration with a large sample size in Telugu speaking adults.

Fourth, significant differences were observed between normophonic and dysphonic individuals for AVQI and five of its six constituent parameters (CPPS, HNR, Shimmer Local %, Shimmer Local dB, and Tilt through LTAS). Dysphonic individuals showed significantly higher AVQI values than normophonic individuals, indicating poorer overall voice quality. Except Slope of LTAS, other constituent parameters of AVQI revealed significant difference between normophonic and dysphonic individuals. Abnormality in vocal fold

vibration, in terms of periodicity, glottal closure, or noise component in voice or tremor in voice may all lead to poor voice quality in individuals with dysphonia. Thus, resulted in abnormal values of AVQI and its constituent parameters in dysphonic individuals.

Fifth, the findings of the present study support the usefulness of AVQI (v02.02) as an objective measure for assessing voice quality and distinguishing dysphonic voices from normophonic voices in Telugu-speaking adults.

Implications of the study

- The present study provides preliminary normative reference values for AVQI (version 02.02) in Telugu-speaking adults between 20 and 65 years of age. These reference values may be useful for future clinical and research applications involving Telugu-speaking populations.
- The findings of the present study provide preliminary information on the influence of age and gender on AVQI and its constituent parameters in Telugu-speaking adults. Further research involving larger samples may help clarify these relationships and contribute to the development of appropriate reference values.
- The normative reference data established in the present study may serve as baseline information for future investigations examining voice quality in Telugu-speaking individuals, including studies involving voice disorders and treatment outcomes.

Limitations of the study

- The number of participants in dysphonic group was small ($N = 13$) compared with the normophonic group ($N = 92$).
- Although efforts were made to maintain consistency in the recording procedure, a small portion of recordings was obtained in a different recording mode (stereo), as a result, these samples were excluded during data screening. This reduced the final sample size available for final analysis.
- Although voice recordings were recorded in a quiet environment, it can be challenging to maintain acoustic background exactly the same during every recording procedure. Minor variations in background noise and recording conditions may have influenced some acoustic measurements in the study.

Future directions

- The present study established normative reference values for AVQI (version 02.02) only in Telugu-speaking adults between 20 and 65 years of age. Future studies may investigate AVQI in other age groups, such as children or beyond 65 years of age.
- The findings related to the comparison between normophonic and dysphonic voices were based on a relatively small sample of individuals with dysphonia. Future studies involving a larger sample of individuals with dysphonia may help establish diagnostic cut-off values for AVQI in Telugu-speaking populations and further examine its clinical utility.

- The present study observed significant gender-related differences in AVQI. Future studies with larger samples may further investigate the influence of gender on AVQI and examine the need for gender-specific reference values.
- The present study used AVQI version 02.02. Future research may establish normative and diagnostic reference values for AVQI version 03.01 in Telugu-speaking adults and compare its performance with version 02.02.
- The applicability of AVQI as an objective measure for monitoring treatment outcomes in voice disorders may be investigated in future studies.

REFERENCES

- American Speech-Language-Hearing Association. (2015). *Voice disorders*.
ASHA Practice Portal. <https://www.asha.org/practice-portal/clinical-topics/voice-disorders/>
- André, D. C., Fernandes, P. O., & Teixeira, J. P. (2022). *Statistical analysis of voice parameters in healthy subjects and with vocal pathologies—HNR*.
<http://hdl.handle.net/10198/27761>
- Aronson, A. E., Bless, D. (2011). *Clinical Voice Disorders*. United States:
Thieme Publishers, New York.
- Awan, S. N., & Roy, N. (2006). Toward the development of an
objective index of dysphonia severity: A four-factor
acoustic model. *Clinical Linguistics & Phonetics*, 20(1),
35–49. <https://doi.org/10.1080/02699200400008353>
- Awan, S. N., & Roy, N. (2015). Acoustic voice analysis: Non-invasive and
accessible voice assessment for clinical settings. *Journal of Voice*, 29(2),
165–172.
- Barsties, B., & Maryn, Y. (2012). Der Acoustic Voice Quality Index in
Deutsch. *HNO*, 60(8), 715–720. <https://doi.org/10.1007/s00106-012-24999>
- Barsties, B., & Maryn, Y. (2015a). Acoustic Voice Quality Index: A
multiparameter objective voice assessment measure. *The Laryngoscope*,
125(11), 2421–2427.
- Barsties, B., & Maryn, Y. (2015b). The improvement of internal
consistency of the Acoustic Voice Quality Index. *American*

Journal of Otolaryngology, 36(5), 647–656.

<https://doi.org/10.1016/j.amjoto.2015.04.012>

Barsties, B., & Maryn, Y. (2016). External validation of the Acoustic Voice Quality Index (AVQI) version 03.01. *JAMA Otolaryngology–Head & Neck Surgery*, 142(12), 1159–1164.

Bauser, M., & Drinnan, M. J. (2011). Variation in the correlations between commonly used acoustic measures and perceptual severity ratings in dysphonia. *Journal of Voice*, 25(3), 330–335.

Benoy, J. J. (2017). Acoustic voice quality index (AVQI) and perceptual measures in the Indian population. Master's Dissertation submitted to the University of Mysore, Mysore.

Bhaskararao, P., & Ray, P. (2018). *Phonetics and Phonology of Retroflex Sounds in Dravidian Languages*. *Journal of South Asian Linguistics*, 6(2), 45–60.

Bhuta, T., Patrick, L., & Garnett, J. D. (2004). Perceptual evaluation of voice quality and its correlation with acoustic measurement. *Journal of Voice*, 18(3), 299–304.

Bough, I. D., Heuer, R. J., Sataloff, R. T., Hills, J. R., & Cater, J. R. (1996). Intrasubject variability of objective voice measures. *Journal of Voice*, 10(2), 166–174. [https://doi.org/10.1016/s0892-1997\(96\)80044-0](https://doi.org/10.1016/s0892-1997(96)80044-0)

Brockmann-Bauser, M., & Drinnan, M. J. (2011). Routine acoustic voice analysis: time to think again? *Current Opinion in Otolaryngology & Head & Neck Surgery*, 19(3), 165–170.
<https://doi.org/10.1097/moo.0b013e32834575fe>

- Calvache, C., Solaque, L., Velasco, A., & Peñuela, L. (2023). Biomechanical models to represent vocal physiology: A systematic review. *Journal of Voice*, 37(3), 465.e1–465.e18. <https://doi.org/10.1016/j.jvoice.2021.02.014>
- Carding, P. N., Wilson, J. A., MacKenzie, K., & Deary, I. J. (2009). Measuring voice outcomes: state of the science review. *The Journal of Laryngology & Otology*, 123(8), 823–829. <https://doi.org/10.1017/s0022215109005398>
- Carding, P., Carlson, E., Epstein, R., Mathieson, L., & Shewell, C. (2000). Formal 4'/perceptual evaluation of voice quality in the United Kingdom. *Logopedics Phoniatrics Vocology*, 25(3), 133-138.
- Coyle, S. M., Weinrich, B. D., & Stemple, J. C. (2001). Shifts in the relative prevalence of laryngeal pathology in a Treatment-Seeking population. *Journal of Voice*, 15(3), 424–440. [https://doi.org/10.1016/s0892-1997\(01\)00043-1](https://doi.org/10.1016/s0892-1997(01)00043-1)
- Dejonckere, P., & Lebacqz, J. (1996). Acoustic, perceptual, aerodynamic and anatomical correlations in voice pathology. *ORL*, 58(6), 326–332. <https://doi.org/10.1159/000276864>
- Delgado Hernández, J., Calzada Uriarte, M., Blanco Mendoza, A., Rodríguez Valdés, R., & Maryn, Y. (2018). Acoustic Voice Quality Index version 03.01 threshold for Spanish speakers. *Acta Otorrinolaringológica Española*, 69(4), 201–207.
- Englert, M. A. C., Barbosa, P. A., & Gois Junior, M. A. (2021). Validation of the Acoustic Voice Quality Index version 03.01 for Brazilian Portuguese speakers. *Journal of Voice*, 35(1), 88–96.

- Fantini, M., Assone, S., Rizzo, A., Mozzanica, F., Ginocchio, D., & Schindler, A. (2023). Validation of the Acoustic Voice Quality Index version 03.01 in Italian-speaking adults. *Journal of Voice*, 37(2), 215–224.
- Fujiki, R. B., & Thibeault, S. L. (2024). Voice disorder prevalence and vocal health characteristics in adolescents. *JAMA Otolaryngology–Head & Neck Surgery*, 150(9), 800–810.
<https://doi.org/10.1001/jamaoto.2024.2081>
- Goy, H., Fernandes, D. N., Pichora-Fuller, M. K., & Van Lieshout, P. (2013). Normative Voice Data for Younger and Older Adults. *Journal of Voice*, 27(5), 545–555. <https://doi.org/10.1016/j.jvoice.2013.03.002>
- Hakkesteegt, M. M., Brocaar, M. P., Wieringa, M. H., & Feenstra, L. (2006). Influence of age and gender on the Dysphonia Severity Index. *Folia phoniatrica et logopaedica*, 58(4), 264-273.
- Hakkesteegt, M. M., Brocaar, M. P., Wieringa, M. H., & Feenstra, L. (2007). The relationship between perceptual evaluation and objective multiparametric evaluation of dysphonia severity. *Journal of Voice*, 22(2), 138–145. <https://doi.org/10.1016/j.jvoice.2006.09.010>
- Heman-Ackah, Y. D., Sataloff, R. T., Laureyns, G., Lurie, D., Michael, D. D., Heuer, R., Rubin, A., Eller, R., Chandran, S., Abaza, M., Lyons, K., Divi, V., Lott, J., Johnson, J., & Hillenbrand, J. (2014). Quantifying the Cepstral Peak Prominence, a Measure of Dysphonia. *Journal of Voice*, 28(6), 783–788.
<https://doi.org/10.1016/j.jvoice.2014.05.005>
- Heman-Ackah, Y. D., Sataloff, R. T., Laureyns, G., Lurie, D., Michael, D. D., Heuer, R., Rubin, A., Eller, R., Chandran, S., Abaza, M., Lyons, K., Divi, V., Lott, J., Johnson, J., & Hillenbrand, J. (2014). Quantifying the Cepstral Peak Prominence, a Measure of Dysphonia. *Journal of Voice*, 28(6), 783–788.

<https://doi.org/10.1016/j.jvoice.2014.05.005>

Hippargekar, P., Bhise, S., Kothule, S., & Shelke, S. (2022). Acoustic Voice Analysis of Normal and Pathological Voices in Indian Population Using Praat Software. *Indian Journal of Otolaryngology and Head & Neck Surgery*, 74(Suppl 3), 5069–5074. <https://doi.org/10.1007/s12070-021-02757-9>

Hirano, M. (1981). *Clinical examination of voice* (Vol. 5). Springer.

Hirano, M., Hibi, S., Yoshida, T., Hirade, Y., Kasuya, H., & Kikuchi, Y. (1988). Acoustic analysis of pathological voice: some results of clinical application. *Acta oto-laryngologica*, 105(5-6), 432-438. <https://doi.org/10.1016/j.jvoice.2006.09.010>

Hosokawa, K., Iwahashi, T., & Kondo, K. (2017). Validation of the Acoustic Voice Quality Index version 02.02 in Japanese speakers. *Journal of Voice*, 31(1), 108–114.

Hosokawa, K., Iwahashi, T., & Kondo, K. (2019). Validation of the Acoustic Voice Quality Index version 03.01 in Japanese speakers. *Journal of Voice*, 33(3), 382–388.

Jayakumar, T., & Benoy, J. J. (2022). Acoustic Voice Quality Index (AVQI) in the Measurement of Voice Quality: A Systematic Review and Meta-Analysis. *Journal of Voice*, 38(5), 1055–1069. <https://doi.org/10.1016/j.jvoice.2022.03.018>

Jayakumar, T., & Savithri, S. (2010). Effect of geographical and ethnic variation on dysphonia Severity Index: A study of Indian population. *Journal of Voice*, 26(1), e11– e16. <https://doi.org/10.1016/j.jvoice.2010.05.008>

- Jayakumar, T., Benoy, J. J., & Yasin, H. M. (2022). Effect of age and gender on acoustic voice Quality Index across the lifespan: a cross-sectional study in the Indian population. *Journal of Voice*, 36(3), 436.e1-436.e8. <https://doi.org/10.1016/j.jvoice.2020.05.025>
- Jayakumar, T., Benoy, J. J., & Yasin, H. M. (2022). Effect of Age and Gender on Acoustic Voice Quality Index Across Lifespan: A Cross-sectional Study in Indian Population. *Journal of Voice*, 36(3), 436.e1-436.e8. <https://doi.org/10.1016/j.jvoice.2020.05.025>
- Jayakumar, T., Rajasudhakar, R., & Benoy, J. J. (2022). Comparison and Validation of Acoustic Voice Quality Index Version 2 and Version 3 among South Indian Population. *Journal of Voice*. <https://doi.org/10.1016/j.jvoice.2022.02.019>
- Kankare, E., Latoszek, B. B. V., Maryn, Y., Asikainen, M., Rorarius, E., Vilpas, S., Ilomäki, I., Tyrmi, J., Rantala, L., & Laukkanen, A. (2020). The acoustic voice quality index version 02.02 in the Finnish-speaking population. *Logopedics Phoniatrics Vocology*, 45(2), 49–56. <https://doi.org/10.1080/14015439.2018.1556332>
- Karnell, M. P. (1991). Laryngeal perturbation analysis: minimum length of analysis window. *Journal of Speech, Language, and Hearing Research*, 34(3), 544- 548.
- Karnell, M. P., Hall, K. D., & Landahl, K. L. (1995). Comparison of fundamental frequency and perturbation measurements among three analysis systems. *Journal of Voice*, 9(4), 383–393. [https://doi.org/10.1016/s0892-1997\(05\)80200-0](https://doi.org/10.1016/s0892-1997(05)80200-0)
- Kempster, G. B., Gerratt, B. R., Abbott, K. V., Barkmeier-Kraemer, J., & Hillman, R. E. (2009). Consensus auditory-perceptual evaluation of

- voice: development of a standardized clinical protocol. *American Journal of Speech-Language Pathology*, 18(2), 124-132.
- Kent, S. A. K., Fletcher, T. L., Morgan, A., Morton, M. E., Hall, R. J., & Sandage, M. J. (2025). Updated Acoustic Normative Data through the Lifespan: A Scoping Review. *Journal of Voice*, 39(4), 1130.e1-1130.e18. <https://doi.org/10.1016/j.jvoice.2023.02.011>
- Kim, G., Lee, Y., Bae, I., Park, H., & Kwon, S. (2018). Application of the New Version of the Acoustic Voice Quality Index with Korean Speakers. *Eon'eo Cheong'gag Jang'ae Yeon'gu/CommunicationSciences &Disorders*, 23(4),1091–1101. <https://doi.org/10.12963/csd.18556>
- Kim, G., Lee, Y., Bae, I., Park, H., Wang, S., & Kwon, S. (2018). Validation of the Acoustic Voice Quality Index in the Korean language. *Journal of Voice*, 33(6), 948.e1- 948.e9. <https://doi.org/10.1016/j.jvoice.2018.06.007>
- Kim, G., Von Latoszek, B. B., & Lee, Y. (2020). Validation of Acoustic Voice Quality Index Version 3.01 and Acoustic Breathiness Index in Korean population. *Journal of Voice*, 35(4), 660.e9-660.e18. <https://doi.org/10.1016/j.jvoice.2019.10.005>
- Kreiman, J., & Gerratt, B. R. (2010). Perceptual assessment of voice quality: Past, present, and future. *Perspectives on Voice and Voice Disorders*, 20, 62-67
- Kreiman, J., Gerratt, B. R., & Berke, G. S. (1994). The multidimensional nature of pathological vocal quality. *Journal of the Acoustical Society of America*, 96, 1291-1302.

Kreiman, J., Gerratt, B. R., Garellek, M., Samlan, R., & Zhang, Z. (2014).

Toward a unified theory of voice production and perception.

Loquens, 1(1).

Latoszek, B. B. V., Lehnert, B., & Janotte, B. (2018). Validation of the Acoustic

Voice Quality Index version 03.01 and Acoustic Breathiness Index in

German. *Journal of Voice*, 34(1), 157.e17-157.e25.

<https://doi.org/10.1016/j.jvoice.2018.07.026>

Laver, J. (1980). The phonetic description of voice quality. Cambridge

Studies in Linguistics London, 31, 1-186.

Laver, J. D., Wirz, S. L., Mackenzie, J., & Miller, S. (1981). A perceptual

protocol for the analysis of vocal profiles. University of Edinburgh

Work in Progress. Linguistics Department, 14.

Lechien, J. R., Geneid, A., Bohlender, J. E., Cantarella, G., Avellaneda, J. C.,

Desuter, G., Sjogren, E. V., Finck, C., Hans, S., Hess, M., Oguz, H.,

Remacle, M. J., Schneider-Stickler, B., Tedla, M., Schindler, A.,

Vilaseca, I., Zabrodsky, M., Dikkers, F. G., & Crevier-Buchman, L.

(2023). Consensus for voice quality assessment in clinical practice:

guidelines of the European Laryngological Society and Union of the

European Phoniaticians. *European Archives of Oto-Rhino-*

Laryngology, 280(12), 5459–5473. [https://doi.org/10.1007/s00405-023-](https://doi.org/10.1007/s00405-023-08211-6)

[08211-6](https://doi.org/10.1007/s00405-023-08211-6)

Liu, H., Wan, Z., Chen, G., & Yin, S. (2025). Correlation between CAPE-V and

GRBAS perceptual scales in dysphonia assessment: A clinical validation

study. *Journal of Voice*, 39(1), 45–52.

<https://doi.org/10.1016/j.jvoice.2023.05.009>

- Maryn, Y., Corthals, P., Van Cauwenberge, P., Roy, N., & De Bodt, M. (2010). Toward improved ecological validity in the acoustic measurement of overall voice quality: Combining continuous speech and sustained vowels. *Journal of Voice*, 24(5), 540–555. <https://doi.org/10.1016/j.jvoice.2008.12.014>
- Maryn, Y., De Bodt, M., Barsties, B., & Roy, N. (2014). The value of the Acoustic Voice Quality Index as a measure of dysphonia severity in subjects speaking different languages. *European Archives of Oto-Rhino-Laryngology*, 271(6), 1609–1619.
- Maryn, Y., Kim, H., & Kim, J. (2016). Auditory-Perceptual and acoustic methods in measuring dysphonia severity of Korean speech. *Journal of Voice*, 30(5), 587–594. <https://doi.org/10.1016/j.jvoice.2015.06.011>
- Maryn, Y., Roy, N., De Bodt, M., Van Cauwenberge, P., & Corthals, P. (2009). Acoustic measurement of overall voice quality: A meta-analysis. *the Journal of the Acoustical Society of American the Journal of the Acoustical Society of America*, 126(5), 2619–2634. <https://doi.org/10.1121/1.3224706>
- Neelanjana, V. (2011). Dysphonia Severity Index in the Indian population. Master's Dissertation, University of Mysore, Mysore, India.
- Nemr, K., Amar, A., Abrahao, M., Leite, G. C., Kohle, J., Santos, A. D., & Correa, L.(2005).Comparative analysis of perceptual evaluation, acoustic analysis and indirect laryngoscopy for vocal assessment of a population with vocal complaint. *Brazilian Journal of Otorhinolaryngology*,71(1), 13-17.

- Nemr, K., Simoes-Zenari, M., Cordeiro, G. F., Tsuji, D., Ogawa, A. I., Ubrig, M. T., & Menezes, M. H. M. (2012). GRBAS and CAPE-V scales: High reliability and consensus when applied at different times. *Journal of Voice*, 26(6), 812.e17–812.e22.
<https://doi.org/10.1016/j.jvoice.2012.03.005>
- Nerurkar, N. K. (Eds.). (2017). *Textbook of Laryngology*. New Delhi: The Health Sciences Publisher.
- Oates, J. (2009). Auditory-Perceptual evaluation of disordered voice quality. *Folia Phoniatrica Et Logopaedica*, 61(1), 49–56.
<https://doi.org/10.1159/000200768>
- Oliveira Santos, A., Godoy, J., Silverio, K., & Brasolotto, A. (2023). Vocal Changes of Men and Women from Different Age Decades: An Analysis from 30 Years of Age. *Journal of Voice*, 37(6), 840–850.
<https://doi.org/10.1016/j.jvoice.2021.06.003>
- Patel, R. R., Awan, S. N., Barkmeier-Kraemer, J., Courey, M., Deliyski, D., Eadie, T., & Hillman, R. (2018). Recommended protocols for instrumental assessment of voice: American Speech-Language-Hearing Association expert panel to develop a protocol for instrumental assessment of vocal function. *American journal of speech-language pathology*, 27(3), 887-905.
- Pebbili, G. K., Shabnam, S., Pushpavathi, M., Rashmi, J., Sankar, R. G., Nethra, R., Shreya, S., & Shashish, G. (2021). Diagnostic Accuracy of Acoustic Voice Quality Index Version 02.03 in Discriminating across the Perceptual Degrees of Dysphonia Severity in Kannada Language. *Journal of Voice*, 35(1), 159.e11-159.e18.
<https://doi.org/10.1016/j.jvoice.2019.07.010>

- Peterson, E. A., Roy, N., Awan, S. N., Merrill, R. M., Banks, R., & Tanner, K. (2013). Toward validation of the Cepstral Spectral Index of Dysphonia (CSID) as an objective treatment outcomes measure. *Journal of Voice*, 27(4), 401–410. <https://doi.org/10.1016/j.jvoice.2013.04.002>
- Piccirillo, J. F., Painter, C., Haiduk, A., Fuller, D., & Fredrickson, J. M. (1998). Assessment of two objective voice function indices. *Annals of Otology Rhinology & Laryngology*, 107(5), 396–400. <https://doi.org/10.1177/000348949810700506>
- Pommée T, Maryn Y, Finck C, et al. The acoustic voice quality index, version 03.01, in French and the voice handicap index. *J Voice*. 2020;34:646.e1–646.e10. <https://doi.org/10.1016/j.jvoice.2018.11.017>
- Reynolds, V., Buckland, A., Bailey, J., Lipscombe, J., Nathan, E., Vijayasekaran, S., Kelly, R., Maryn, Y., & French, N. (2012). Objective assessment of pediatric voice disorders with the acoustic voice quality index. *Journal of Voice*, 26(5), 672-e1-672.e7.
- Rojas, S., Kefalianos, E., & Vogel, A. (2020). How Does Our Voice Change as We Age? A Systematic Review and Meta-Analysis of Acoustic and Perceptual Voice Data From Healthy Adults Over 50 Years of Age. *Journal of Speech, Language, and Hearing Research*, 63(2), 533–551. https://doi.org/10.1044/2019_JSLHR-19-00099
- Roy, N., Merrill, R. M., Gray, S. D., & Smith, E. M. (2005). Voice Disorders in the general Population: Prevalence, risk factors, and occupational impact. *The Laryngoscope*, 115(11), 1988–1995. <https://doi.org/10.1097/01.mlg.0000179174.32345.41>

- Savithri, S.R and Jayaram, M. (2005).Rate of Speech/Reading:Project Report(Sanction No.SH/ARF-RSRDL-17/2004-2005)All India Institute of Speech and Hearing,Mysuru
- Scherer, R. C., Vail, V. J., & Guo, C. G. (1995). Required number of tokens to determine representative voice perturbation values. *Journal of Speech Language and Hearing Research*, 38(6), 1260–1269. <https://doi.org/10.1044/jshr.3806.1260>
- Seshasri, D. (2018). *Acoustic voice quality index in Kannada speaking children in the age range of 10 to 12 years*. Masters Dissertation submitted to the University of Mysore, Mysuru.
- Shabnam, S., & Pushpavathi, M. (2021). Effect of gender on Acoustic Voice Quality Index02.03 and dysphonia severity Index in Indian normophonic adults. *Indian Journal of Otolaryngology and Head & Neck Surgery*, 74(S3), 5052–5059. <https://doi.org/10.1007/s12070-021-027128>
- Shabnam, S., & Pushpavathi, M. (2022). Effect of Gender on Acoustic Voice Quality Index 02.03 and Dysphonia Severity Index in Indian Normophonic Adults. *Indian Journal of Otolaryngology and Head & Neck Surgery*, 74(Suppl 3), 5052–5059. <https://doi.org/10.1007/s12070-021-02712-8>
- Teixeira, J. P., & Fernandes, P. O. (2014). Jitter, Shimmer and HNR Classification within Gender, Tones and Vowels in Healthy Voices. *Procedia Technology, CENTERIS 2014 - Conference on ENTERprise Information Systems / ProjMAN 2014 - International Conference on Project MANagement / HCIST 2014 - International Conference on Health and Social Care Information Systems and Technologies*, 16, 1228–1237. <https://doi.org/10.1016/j.protcy.2014.10.138>

- Thomson, S. L. (2024). Synthetic, self-oscillating vocal fold models for voice production research. *The Journal of the Acoustical Society of America*, 156(2), 1283–1308. <https://doi.org/10.1121/10.0028267>
- Timmermans, B., De Bodt, M. S., Wuyts, F. L., & Van De Heyning, P. H. (2004). Training Outcome in Future Professional Voice Users after 18 Months of Voice Training. *Folia Phoniatrica Et Logopaedica*, 56(2), 120–129. <https://doi.org/10.1159/000076063>
- Uloza, V., Latoszek, B. B. V., Ulozaite-Staniene, N., Petrauskas, T., & Maryn, Y. (2018). A comparison of Dysphonia Severity Index and Acoustic Voice Quality Index measures in differentiating normal and dysphonic voices. *European Archives of Oto-Rhino- Laryngology*, 275(4), 949–958. <https://doi.org/10.1007/s00405-018-4903-x>
- Vallur, M., & Jagadeeswaran, M. (2025). Correlation between GRBAS perceptual ratings and patient-reported outcomes in voice disorders: A systematic review. *Journal of Voice*. Advance online publication. <https://doi.org/10.1016/j.jvoice.2024.10.021>
- Varna, K S.(2024). Acoustic Voice Quality Index (AVQI) for Kannada and Malayalam speaking normophonic children.
- Vaz Freitas, S., Pestana, P. M., Almeida, V., & Ferreira, A. (2014). Audio-perceptual evaluation of Portuguese voice disorders—An inter- and intrajudge reliability study. *Journal of Voice*, 28(2), 210–215. <https://doi.org/10.1016/j.jvoice.2013.07.010>

- Vishali, P. (2019). Acoustic Voice Quality Index (AVQI) in Tamil Language. Master's Dissertation submitted to the University of Mysore, Mysore.
- Webb, A.L., Carding, P.N., Deary, I.J. *et al.* The reliability of three perceptual evaluation scales for dysphonia. *Eur Arch Otorhinolaryngol* **261**,429–434(2004).
<https://doi.org/10.1007/s00405-003-0707-7>
- Wilson, D. K. (1987). Voice problems of children (3rd ed.). Williams & Wilkins.
- Wolfe, V., Fitch, J., & Cornell, R. (1995). Acoustic prediction of severity in commonly occurring voice problems. *Journal of Speech Language and Hearing Research*, **38**(2), 273–279.
<https://doi.org/10.1044/jshr.3802.273>
- Wuyts, F. L., De Bodt, M. S., Molenberghs, G., Remacle, M., Heylen, L., Millet, B., Van Lierde, K., Raes, J., & Van De Heyning, P. H. (2000). The Dysphonia Severity Index. *Journal of Speech Language and Hearing Research*, **43**(3), 796– 809.
<https://doi.org/10.1044/jslhr.4303>
- Yamauchi, A., Yokonishi, H., Imagawa, H., Sakakibara, K.-I., Nito, T., Tayama, N., & Yamasoba, T. (2015). Quantitative Analysis of Digital Videokymography: A Preliminary Study on Age- and Gender-Related Difference of Vocal Fold Vibration in Normal Speakers. *Journal of Voice*, **29**(1), 109–119.
<https://doi.org/10.1016/j.jvoice.2014.05.006>
- Yu, P., Ouaknine, M., Revis, J., & Giovanni, A. (2001). Objective voice analysis for dysphonic patients: a multiparametric

protocol including acoustic and aerodynamic measurements. *Journal of voice*, 15(4), 529-542.

Zainae, S., Khadivi, E., Jamali, J., Sobhani-Rad, D., Maryn, Y., & Ghaemi, H. (2022). The acoustic voice quality index, version 2.06 and 3.01, for the Persian-speaking population. *Journal of Communication Disorders*, 100, 106279.

<https://doi.org/10.1016/j.jcomdis.2022.106279>

Ziwei, Y., Zheng, P., & Pin, D. (2014). Multiparameter voice assessment for voice disnguage *Pathology*, 20(1), 14-22.

Zraick, R. I., Kempster, G. B., Connor, N. P., Thibeault, S., Klaben, B. K., Bursac, & Glaze, L. E. (2011). Establishing validity of the consensus auditory- perceptual evaluation of voice (CAPE-V). *American Journal of Speech- Language Pathology*, 20(1), 14-22.

APPENDIX - A

Informed Consent

I have been informed about the procedures of the study. The possible risks and benefits of my participation as a human subject in the study are clearly understood by me. I understand that I have a right to refuse participation as a subject or withdraw my consent at any time. I am also aware that by subjecting this investigation, I must give more time for assessments by the investigating team and that these assessments may/may not benefit me. I have the freedom to write to Chairman AEC, in case of any violation of these provisions without the danger of me being denied any rights to secure the clinical services at this institute.

I, _____ the undersigned, give my consent to be a participant in this investigation/study/program.

Signature of Participant

(Name and address)

Signature of the researcher:

Name and Designation: Ms. Keerthana Katti, II MSc. SLP

Date:

APPENDIX B

Reading Passage in Telugu (Source: Savithri & Jayaram, 2005)

ధనపతి అనే వ్యాపారి ఉండేవాడు, బొజ్జ బ్రాహ్మణునికి ఇతనికి బాల్యం నుంచే స్నేహం ఉండేది. ఒకసారి వ్యాపారి కుమారునికి పుట్టినరోజు పండగ జరిగింది. ఆ కార్యక్రమానికి అనేకమంది పెద్దలను ఆహ్వానించాడు. బ్రాహ్మణుణ్ణి కూడా ఆహ్వానించాడు. ప్రీయమైన స్నేహితుని ఇల్లు, దానికి తోడు షావుకారు స్నేహితుని ఇంటి కార్యక్రమం, తలచుకుంటేనే నోటినిండా నీళ్ళూరాయి.

APPENDIX C

ETHICAL CLEARANCE

**All India Institute of Speech and Hearing**

(An autonomous Institute under the
Ministry of Health and Family Welfare, Govt. of India)
Center of Excellence - Assessed & Accredited by NAAC with 'A' Grade
ISO 9001:2015 Implemented Institute
Manasagangothri, Mysuru-570006

ಅಖಿಲ ಭಾರತ ವಾಕ್ ಮತ್ತು ಶ್ರವಣ ಸಂಸ್ಥೆ
ಮಾನಸಗಂಗೋತ್ರಿ, ಮೈಸೂರು-570006
अखिल भारतीय वाक् श्रवण संस्थान
मानसगंगोत्री, मैसूरु - 570 006

**INSTITUTE REVIEW BOARD (IRB) APPROVAL FOR BIO-BEHAVIOURAL
RESEARCH PROPOSAL INVOLVING HUMAN SUBJECTS AT AIISH**

AIISH Institute Review Board (IRB)

Ref: SH/IRB/M.5/SLP.21/2025-26

Date: 11.03.2026

Title of the Dissertation

: Establishing Normative Reference
Data for AVQI in Telugu Speaking
Adults

Name of the Student

: Ms. Keerthana Katti

Guide

: Dr. R. Rajasudhakar

Co-Guide (if any)

: NIL

Proposed Duration of the Project

: 06 months

Source of funding

: Non-funded Master's dissertation

Estimated budget

: Nil

Reference Number of the Project

: Nil

Date on which the IRB meeting was held

: 05.03.2026

**A clear statement of the decision reached IRB meeting
(In the event of the proposal is not APPROVED,
a statement of reasons for the same must be indicated)**

: APPROVED

Advice & suggestions (if any)

: Advised to submit the
permission letters from clinics
where the participants are being
recruited for the study.

Signature & Name of the Member Secretary-IRB

Dr. Animesh Barman
Professor of Audiology

All India Institute of Speech and Hearing, Mysore-570006

अखिल भारतीय वाक् श्रवण संस्थान

All India Institute of Speech and Hearing

मानसगंगोत्री / Manasagangothri

Phone : 0821-2502000 / 2502100 Toll Free No. 18004255218 Fax : 0821-2510515

e-mail : director@aiishmysore.in / @aiishmysuruofficial, @aiishmysuru

website : www.aiishmysore.in